

How Eating Meat Contributes to Global Warming (*page 72*)

SCIENTIFIC AMERICAN

Nanomedicine

Revolutionizing
the Fight
against Disease

page 44

February 2009

www.SciAm.com



Black holes may have even
stranger siblings that violate
known laws of physics

NAKED SINGULARITIES



Fold the Brain

How Forces That Shape It
Link to Autism and More

Submarine Lava

The Mysterious, Red-Hot
Origins of the Ocean Floor

Electric Rockets

Efficient Travel to Deep Space



IN REALITY, THERE'S NO SUCH THING AS CLEAN COAL.



Coal is one of the leading causes of global warming. But that hasn't stopped the coal industry from advertising clean coal. Yet, the truth is there isn't a single commercial coal power plant in America today that captures its global warming pollution. Learn more about what the coal industry is not telling you at thisisreality.org



Launching into a new world of ultra high-energy gamma rays and deep-space physics...

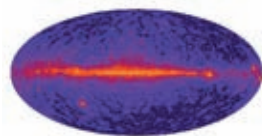


In June of 2008, this Delta II rocket launched the new Fermi Gamma-ray Space Telescope into an orbit 350 miles above the Earth.

The Fermi Gamma-ray Space Telescope

Orbiting the Earth, the new Fermi Telescope scans the entire sky every three hours, collecting gamma-ray information from deep space. And Hamamatsu photonics components are assisting...

In the *Tracker* module of the *Large Area Telescope* (LAT) our silicon strip sensors help to precisely measure angles of incoming gamma-ray particles.



A full-sky scan from the Fermi

In the LAT *Calorimeter*, which measures particle energies, our silicon photodiodes are at work.

And our photomultiplier tubes (PMTs) are incorporated into the LAT *Anticoincidence Detector* (ACD), which filters out cosmic rays.

Hamamatsu ruggedized PMTs are also used in

Hamamatsu is opening the new frontiers of Light * * *

the special *Gamma-ray Burst Monitor* (GBM) to help detect very low-energy gamma rays.

Opening new frontiers

The Fermi is the first gamma-ray observatory that is able to survey the entire sky every day with high sensitivity. And already it has discovered a unique new pulsar that "blinks" gamma rays!



Silicon strip sensor (upper) and PMT sensor

The Fermi Telescope will be a powerful tool for learning more about the early universe, studying phenomena such as black holes and unlocking the mysteries of high-energy particles in deep space. Hamamatsu is proud to be helping!

<http://jp.hamamatsu.com/en/rd/publication/>

HAMAMATSU

The Frontiers of Light

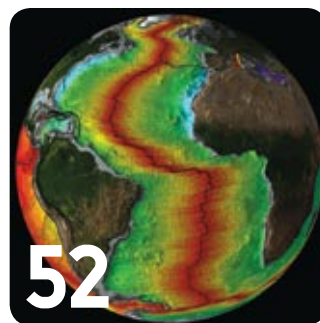
36 **PHYSICS** Naked Singularities

By Pankaj S. Joshi

The black hole has a troublesome sibling, the naked singularity. It is a point where space and time break down. Physicists have long thought—hoped—it could never exist. But they might be wrong.



44



52



58



ON THE COVER

If singularities can exist outside the protective envelope of an event horizon, they may distort the universe in unimaginable ways.

Image by Kenn Brown, Mondolithic Studios.

44 **MEDICINE** Nanomedicine Targets Cancer

By James R. Heath, Mark E. Davis and Leroy Hood

By viewing the body as a system of molecular networks, future physicians will be able to target disruptions in that system with nanoscale technologies and thereby transform the diagnosis and treatment of malignancies and other diseases.

52 **EARTH SCIENCE** The Origin of the Land under the Sea

By Peter B. Kelemen

The deep basins under the oceans are carpeted with lava that spewed from submarine volcanoes—in a most unusual way.

58 **SPACE TECHNOLOGY** New Dawn for Electric Rockets

By Edgar Y. Choueiri

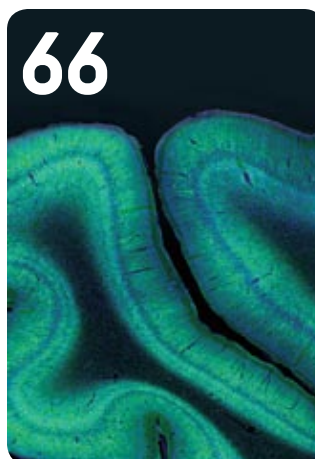
Efficient electric plasma engines are already propelling the next generation of space probes to the outer solar system. That is just the beginning.

NEUROSCIENCE

66 Sculpting the Brain

By Claus C. Hilgetag and Helen Barbas

Knowing how the brain's convolutions take shape could aid the diagnosis and treatment of autism, schizophrenia and other mental disorders.



CLIMATE CHANGE

72 The Greenhouse Hamburger

By Nathan Fiala

Producing beef for the table has a surprising environmental cost: it releases prodigious amounts of heat-trapping greenhouse gases.



GO TO SCIAM.COM

IN-DEPTH REPORT: THE 2009 CONSUMER ELECTRONICS SHOW ▼

On-the-ground reports from CES, the annual technology and gadget fiesta in Las Vegas.

More at www.SciAm.com/feb2009



COURTESY OF EINZIG/INTERNATIONAL CES

Science Talk

The Science of Pain

Stanford University expert Sean Mackey explores the modern understanding of pain, why therapy is so important, and the relation between pain and empathy.

60-Second Science Blog

Dispatches from the Bottom of the Earth

Marine geophysicist Robin Bell reports from Antarctica as she leads an expedition to explore a surprising mountain range underneath the ice sheet.

Extreme Tech

Resuscitating the Atomic Airplane:

Flying on a Wing and an Isotope

Could nuclear-powered planes help save the environment? Engineers reconsider a cold war-era proposal scrapped decades ago.

60-Second Psych Podcast

Cyberchondria: Online Diagnosis Leads to Obsessive Fear

Beware using the Web for self-diagnosis: you will probably end up with a lot of unnecessary stress.

News

Rock and Roll: Meteorites Hitting Early Earth's Oceans May Have Helped Spawn Life

Did heat, pressure and carbon from meteorite impacts create biological precursors?

Scientific American (ISSN 0036-8733), published monthly by Scientific American, Inc., 415 Madison Avenue, New York, N.Y. 10017-1111. Copyright © 2009 by Scientific American, Inc. All rights reserved. No part of this issue may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying and recording for public or private use, or by any information storage or retrieval system, without the prior written permission of the publisher. Periodicals postage paid at New York, N.Y., and at additional mailing offices. Canada Post International Publications Mail (Canadian Distribution) Sales Agreement No. 40012504. Canadian BN No. 127387652RT; QST No. Q1015332537. Publication Mail Agreement #40012504. Return undeliverable mail to Scientific American, P.O. Box 819, Stn Main, Markham, ON L3P 8A2. Subscription rates: one year \$34.97, Canada \$49 USD, International \$55 USD. Postmaster: Send address changes to Scientific American, Box 3187, Harlan, Iowa 51537. Reprints available: write Reprint Department, Scientific American, Inc., 415 Madison Avenue, New York, N.Y. 10017-1111; (212) 451-8877; fax: (212) 355-0408. Subscription inquiries: U.S. and Canada (800) 333-1199; other (515) 248-7684. Send e-mail to sacust@sciam.com Printed in U.S.A.





KUWAIT PRIZE 2009

Invitation for Nominations

The **Kuwait Foundation for the Advancement of Sciences (KFAS)** institutionalized the **KUWAIT Prize** to recognize distinguished accomplishments in the arts, humanities and sciences. The Prizes are awarded annually in the following categories:

- A. Basic Sciences
- B. Applied Sciences
- C. Economics and Social Sciences
- D. Arts and Literature
- E. Arabic and Islamic Scientific Heritage

The Prizes for **2009** will be awarded in the following fields:

- | | |
|---|---|
| 1. Basic Sciences | • Physics |
| 2. Applied Sciences | • Cancer Diseases |
| 3. Economic and Social Sciences | • Privatization Programs and their Effects on Development in the Arab World |
| 4. Arts and Literature | • Studies in Children Literature |
| 5. Arabic and Islamic Scientific Heritage | • City Planning and Topography |

Foreground and Conditions of the Prize:

1. Two prizes are awarded in each category:
 - * A Prize to recognize the distinguished scientific research of a Kuwaiti citizen, and,
 - * A Prize to recognize the distinguished scientific research of an Arab citizen.
2. The candidate should not have been awarded a Prize for the submitted work by any other institution.
3. Nominations for these Prizes are accepted from individuals, academic and scientific centers, learned societies, past recipients of the Prize, and peers of the nominees. No nominations are accepted from political entities.
4. The scientific research submitted must have been published during the last ten years.
5. Each Prize consists of a cash sum of K.D. 30,000/- (approx. U.S.\$100,000/-), a Gold medal, a KFAS Shield and a Certificate of Recognition.
6. Nominators must clearly indicate the distinguished work that qualifies their candidate for consideration.
7. The results of KFAS decision regarding selection of winners are final.
8. The documents submitted for nominations will not be returned regardless of the outcome of the decision.
9. Each winner is expected to deliver a lecture concerning the contribution for which he was awarded the Prize.

Inquiries concerning the KUWAIT PRIZE and nominations including complete curriculum vitae and updated lists of publications by the candidate with **four copies** of each of the published papers should be received before **31/10/2009** and addressed to:

The Director General

The Kuwait Foundation for the Advancement of Sciences - P.O. Box: 25263, Safat - 13113, Kuwait.

Tel: (+965) 22429780 / Fax: 22403891 / E-Mail: prize@kfasc.org.kw

8 From the Editor

10 Letters

14 50, 100 and 150 Years Ago

16 Updates

18 NEWS SCAN

- Cancer drugs for the poorest.
- Faster crosstown traffic.
- Giant gerbils to fight plague.
- Preparing for a quantum afterlife.
- How the turtle got its shell.
- “Lazy eye” and brain rejuvenation.
- Acid trip for the oceans.
- Data Points: Shorter lives for zoo elephants.

OPINION

32 ■ SciAm Perspectives

A scientific stimulus for the U.S.

33 ■ Sustainable Developments

By Jeffrey D. Sachs

Transforming the auto industry.

34 ■ Skeptic

By Michael Shermer

Two myths persist about evolution and natural selection.

35 ■ Anti Gravity

By Steve Mirsky

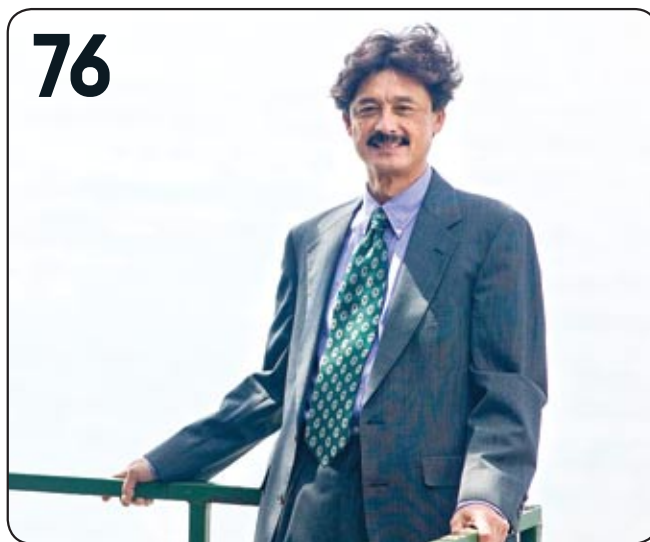
Does airport security keep us unkempt?

14



16

76



GARY PAYNE

76 Insights

Mathematician George Sugihara argues that complexity theory offers a counterintuitive way to revitalize the world's shrinking fisheries.

80 Working Knowledge

Touch screens for telephones.

82 Reviews

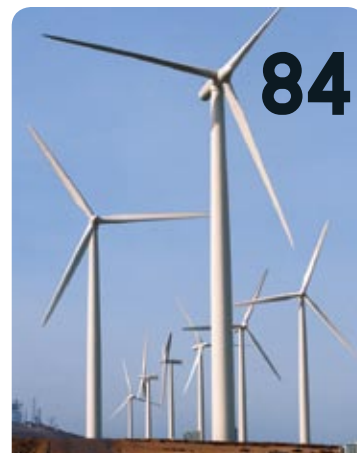
The fourth dimension. Extraterrestrial life. Happy Birthday, Charles Darwin.

84 Ask the Experts

Why do wind turbines have three blades, and my fan at home has five?

What happens to the donor's DNA in a blood transfusion?

84



32

34



Our natural resources are worth protecting. That's why Waste Management works with communities and the Wildlife Habitat Council to use the property adjacent to our landfills as safe havens for native animal and plant life. You might say it's in our nature to do what's good for the environment.

There's more to think about at **ThinkGreen.com**

*From everyday collection
to environmental protection, Think Green.[®]
Think Waste Management.*



A Molecular Checkup

Nanomedicine is coming, but Raquel Welch (alas!) need not apply



Not long ago cancer medicine in the U.S. passed a hopeful milestone: for the first time, the incidence rates for both new cases and deaths in men

and women declined, according to an annual report issued in late November from the National Cancer Institute, the American Cancer Society and other leading organizations. Between 1999 and 2005 diagnosis rates dropped annually by about 0.8 percent. Although deaths from some specific conditions have gone up, overall mortality from cancer is on the decline for both men and women of almost all ethnic groups, as it has been since the early 1990s, in large part because of a shrinking toll from malignancies of the lung, prostate, breast and colon.

That good news invites some cautious interpretation. Incidence rates might have fallen because fewer patients are going for mammograms, prostate screening tests and other diagnostic procedures; if so, physicians may not yet be aware of cases that will eventually surface. The drop in the mortality statistics may largely reflect the population's healthier way of life—most significantly, its decision to kick the tobacco habit. That development is highly welcome, but it may be hard to maintain as a trend: How many other changes can people make that will be so beneficial?

To keep this anticancer momentum, therefore, health care will surely need to step up prevention and treatment in ways that are more tolerable (and affordable) for the general public. The evolving clinical field of nanomedicine may hold many of the answers, as biomedical researchers James R. Heath, Mark E. Davis and Leroy Hood describe in their article, beginning on page 44.

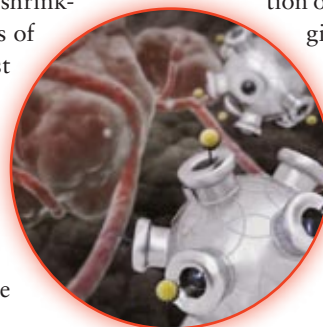
The term “nanomedicine” still conjures up images of teams of microscopic robots performing lifesaving surgery inside our tissues, like the miniaturized submarine crew in the 1966 movie *Fantastic Voyage*. Given the state of the relevant technologies, any such possibility seems at least decades away. (Personally, I will have more confidence in nanobot surgical teams sometime well after engineers can build, say, crews of autonomous dog-size robots that can keep bridges and tunnels in good repair.) But this does not mean nanomedicine is as vaporous.

Rather, in the same way that nanotechnology is better understood as the application of quantum mechanics to engineering, not the use of atoms as building blocks, nanomedicine might best be viewed as a systemic approach to understanding and maintaining health at the molecular level. As Heath, Davis and Hood explain, the accelerating advance of genomic science makes it easier to identify the hallmarks of illnesses even when no

symptoms may be apparent to the patient or clinician. Not only could new blood tests diagnose a nascent liver tumor, for example, but they could also determine to which subcategory of brain tumor it belonged and suggest which gene-focused treatments might be most effective.

Such measures could not always guarantee a cure, but they might someday be able to make often fatal diseases such as cancer and AIDS manageable, much as diabetes is now. The brilliant paradox of nanomedicine is that by focusing on what is extremely small, it can provide a better way to treat a whole person.

JOHN RENNIE
editor in chief



TARGETED NANOPARTICLES could attack tumor cells selectively.

Among Our Contributors



HELEN BARBAS

is a professor at Boston University, where her research focuses on the circuit patterns of the prefrontal cortex of the brain.



EDGAR Y. CHOEIRI

directs the Electric Propulsion and Plasma Dynamics Laboratory at Princeton University, where he also teaches astronautics and applied physics.



NATHAN FIALA

is a doctoral candidate in economics at the University of California, Irvine. His study of the environmental impact of meat production was recently published in *Ecological Economics*.



CLAUS C. HILGETAG

is a founding member of the faculty at Jacobs University Bremen in Germany, where he serves as associate professor of neuroscience.



PANKAJ S. JOSHI

is professor of physics at the Tata Institute of Fundamental Research in Mumbai, where he investigates gravitation and cosmology.



IN THE NEW ENERGY FUTURE IF IT DOESN'T EXIST WE'LL NEED TO INVENT IT.

Tackling climate change and providing fuel for a growing population seems like an impossible problem, but at Shell we try to think creatively.

In addition to our growing oil and gas businesses, we're investing in energy sources like wind, and also investigating innovative new engine fuels made from unexpected sources like gas, hydrogen, straw, waste woodchips and marine algae.

It won't be easy. Innovative solutions rarely are. But, when the challenge is hardest, when everyone else is shaking their heads, we believe there is a way.

To find out how Shell is helping prepare for the new energy future visit **www.shell.com/us/realenergy**



SCIENTIFIC AMERICAN®

Established 1845

EDITOR IN CHIEF: John Rennie
EXECUTIVE EDITOR: Mariette DiChristina
MANAGING EDITOR: Ricki L. Rusting
CHIEF NEWS EDITOR: Philip M. Yam
SENIOR WRITER: Gary Stix
EDITORS: Peter Brown, Graham P. Collins, Mark Fischetti, Steve Mirsky, George Musser, Christine Soares, Kate Wong
CONTRIBUTING EDITORS: Mark Alpert, Steven Ashley, Stuart F. Brown, W. Wayt Gibbs, Marguerite Holloway, Christie Nicholson, Michelle Press, Michael Shermer, Sarah Simpson

MANAGING EDITOR, ONLINE: Ivan Oransky
NEWS EDITOR, ONLINE: Lisa Stein
ASSOCIATE EDITORS, ONLINE: David Biello, Larry Greenemeier
NEWS REPORTERS, ONLINE: Coco Ballantyne, Jordan Lite, John Matson
ART DIRECTOR, ONLINE: Ryan Reid

ART DIRECTOR: Edward Bell
SENIOR ASSOCIATE ART DIRECTOR: Mark Clemens
ASSISTANT ART DIRECTOR: Johnny Johnson
ASSISTANT ART DIRECTOR: Jen Christiansen
PHOTOGRAPHY EDITOR: Monica Bradley
PRODUCTION EDITOR: Richard Hunt

COPY DIRECTOR: Maria-Christina Keller
COPY CHIEF: Daniel C. Schlenoff
COPY AND RESEARCH: Michael Battaglia, Aaron Shattuck, Rachel Dvoskin, Aaron Fagan, Michelle Wright, Ann Chin

EDITORIAL ADMINISTRATOR: Avonelle Wing
SENIOR SECRETARY: Maya Harty

ASSOCIATE PUBLISHER, PRODUCTION: William Sherman
MANUFACTURING MANAGER: Janet Cermak
ADVERTISING PRODUCTION MANAGER: Carl Cherebin
PREPRESS AND QUALITY MANAGER: Silvia De Santis
PRODUCTION MANAGER: Christina Hippeli
CUSTOM PUBLISHING MANAGER: Madelyn Keyes-Milch

BOARD OF ADVISERS

RITA R. COLWELL
Distinguished Professor, University of Maryland
College Park and Johns Hopkins Bloomberg School
of Public Health

DANNY HILLIS
Co-chairman, Applied Minds

VINOD KHOSLA
Founder, Khosla Ventures

M. GRANGER MORGAN
Professor and Head of Engineering and Public
Policy, Carnegie Mellon University

LISA RANDALL
Professor of Physics, Harvard University

GEORGE M. WHITESIDES
Professor of Chemistry, Harvard University

LETTERS

editors@SciAm.com

Intelligence ■ Loop Quantum Gravity ■ Monty Hall



OCTOBER 2008

■ Inscrutable Intelligence?

In "The Search for Intelligence," Carl Zimmer relates the search for evidence of genetic influences on intelligence by Robert Plomin of the Institute of Psychiatry in London and others. Plomin's many genetic studies have so far found only one plausible candidate for an "intelligence gene," and it explains less than 1 percent of the variance on IQ tests. After such an extensive search, it seems likely that absence of evidence is evidence of absence.

Much of the argument for a genetic influence on intelligence is based on IQ test results, but there is very little theoretical backing for what these are testing. Philosophers such as Keith DeRose of Yale University maintain that whether you even *know* something depends on the context, including the stakes for getting it wrong. Many experimental results agree: Claude M. Steele of Stanford University, for instance, found that African-American students underperformed white students on a test that was framed as being diagnostic of intelligence but performed just as well as whites when this verbiage was absent.

There remains a widespread belief in genetic intelligence differences even though there is so little evidence, which makes the idea of intelligence genes resemble folk psychology. Such ideas do real damage in the classroom and beyond. For example, Carol S. Dweck of Stanford University has shown that students believing in pre-set intelligence perform worse than their counterparts who believe that intelligence

"After such an extensive search for an 'intelligence gene,' it seems likely that absence of evidence is evidence of absence."

—Luke Conlin COLLEGE PARK, MD.

is dynamic and results from hard work.

Given the shaky theoretical and experimental backing as well as the potential harm that belief in intelligence genes can do, I hope that researchers will start to question whether searching for a genetic basis of intelligence is worth the risk.

Luke Conlin
College Park, Md.

■ Bangs, Bounces and Black Holes

In "Follow the Bouncing Universe," Martin Bojowald makes a compelling argument for the theory of loop quantum gravity, in which space is subdivided into "atoms" of volume and has a finite capacity to store matter and energy. This structure would prevent singularities, meaning that our universe may have existed before the big bang. Bojowald suggests that previous to the big bang, the universe may have undergone an implosion that was then reversed in a "big bounce," followed by the big bang. He implies, I think, that this cycle may have been going on eternally. But in our current understanding, dark energy seems to promise that the universe will expand forever. Has our universe experienced its last bounce?

Robert Snyder
Andover, Mass.

According to loop quantum gravity theory, true singularities cannot exist. What, then, becomes of black holes?

Mark Saha
Santa Monica, Calif.

**AlwaysOn**
The insider's network

&

**SCIENTIFIC
AMERICAN**

Co-Present



AlwaysOn GoingGreen^{east}

March 9th–11th, 2009
Four Seasons Hotel, Boston, MA

**EARLY
REGISTRATION
STILL AVAILABLE!
REGISTER TODAY
AND GET A SPECIAL
DISCOUNT.**

PLEASE VISIT
GOINGGREENTICKETS.COM
FOR DETAILS

Greener Pastures for Global Business

GoingGreen East

Like its sister event in San Francisco, GoingGreen East is where cutting-edge greentech CEOs meet top investors and the movers and shakers from the biggest industries on earth. This two-and-a-half-day executive event features CEO presentations and high-level debates on the most promising emerging green technologies and new entrepreneurial opportunities. At GoingGreen East, our editors will honor the GoingGreen Top 50 Private Companies. Up to 40 greentech CEOs will also pitch their market strategies to a panel of industry experts in our "CEO Showcases."

Who Participates in GoingGreen East

Seven hundred greentech CEOs, business development officers, eminent researchers, venture capital and private-equity investors, and leading members of the press and blogging community will attend GoingGreen East. Thousands of webcast viewers from over 100 countries will also tune in and interact with the program. Executives attend GoingGreen East to identify and debate emerging trends, build high-level relationships and create new business opportunities.

Sponsorship information:

Marc Sternberg
President & COO
AlwaysOn
310-403-3330
marc@alwayson-network.com

Bruce Brandfon
Vice President & Publisher
Scientific American
212-451-8561
bbrandfon@sciam.com

Program:

If you have feedback on our program or would like to suggest speakers or panel topics, contact:

Ed Ring
ed@ecoworld.com

Tickets:

If you want to order your ticket by phone or inquire about group rates, contact:

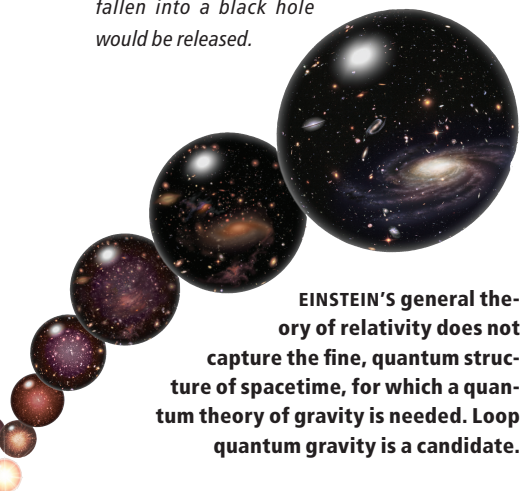
Michelle Hughes
michelle@alwayson-network.com
415.367.9538 ext 103

For more information, visit www.GoingGreenTickets.com

**AlwaysOn**
The insider's network

BOJOWALD REPLIES: *In response to Snyder's question, our big bang may have been the last (or only) bounce. Because we do not really know what dark energy is, however, its existence could well be the result of an intermediate form of matter, which might decay in the future. In this case, there can still be a future bounce. Observations in the not too distant future will tell us more about dark energy.*

Regarding Saha's letter, black holes are more difficult to analyze in quantum gravity than other regions of space because they are not as symmetric. (The gravitational force changes with the distance from the black hole, making the situation inhomogeneous.) Currently there are indications that the center of black holes is indeed nonsingular though still very dense. The usual horizons of black holes would form in loop quantum gravity, but light would be trapped only for a finite period. After the center bounces back, much like in cosmology, what has fallen into a black hole would be released.



EINSTEIN'S general theory of relativity does not capture the fine, quantum structure of spacetime, for which a quantum theory of gravity is needed. Loop quantum gravity is a candidate.

■ Prize Probabilities

In "A Random Walk through Middle Land" [Skeptic], Michael Shermer presents the following scenario: you are a contestant on *Let's Make a Deal* and are shown three doors. Behind one is a new car; the others hide goats. You choose a door, and host Monty Hall reveals a goat behind a different door. Shermer then posits that you have a two-thirds chance of winning by switching your choice because there are only three possible door configurations (good, bad, bad; bad, good, bad; and bad, bad, good), and with the latter two you win by switching. But he fails to recognize that the second configuration has been taken off the table. You have a 50 percent chance.

Andrew Howard
Los Angeles

SHERMER REPLIES: *In nearly 100 months of writing the Skeptic column, I have never received so many letters as I did disagreeing with my description of the so-called Monty Hall Problem. James Madison University mathematics professor Jason Rosenhouse, who has written an entire book on the subject—The Monty Hall Problem: The Remarkable Story of Math's Most Contentious Brain Teaser (Oxford University Press, 2009)—explained to me that you double your chances of winning by switching doors when three conditions are met: 1) Monty never opens the door you chose initially; 2) Monty always opens a door concealing a goat; 3) When your initial choice was correct, Monty chooses a door at random. "Switching turns a loss into a win and a win into a loss," Rosenhouse says. "Since my first choice is wrong two thirds of the time, I will win that often by switching."*

At the beginning you have a one-third chance of picking the car and a two-thirds chance of picking a goat. Switching doors is bad only if you initially chose the car, which happens one third of the time, and switching doors is good if you initially chose a goat, which happens two thirds of the time. Thus, the probability of winning by switching is two thirds. Analogously, if there are 10 doors, initially you have a one-tenth chance of picking the car and a nine-tenths chance of picking a goat. Switching doors is bad only if you initially chose the car, which happens one tenth of the time. So the probability of winning by switching is nine tenths—assuming that Monty has shown you eight other doors with goats.

Still not convinced? Google "Monty Hall Problem simulation" and try the various computer simulations. You will see that you double your actual wins by switching doors. One of my skeptical correspondents ran his own simulation more than 10,000 trials, concluding that "switching doors yields a two-thirds success rate while running without switching doors yields a one-third success rate." (Go to <http://tinyurl.com/bu9jl> for the simulation.)

➔ Read an expanded version of Howard's letter and Shermer's reply at www.SciAm.com/feb2009

Letters to the Editor

Scientific American
415 Madison Ave.
New York, NY 10017-1111
or editors@SciAm.com

Letters may be edited for length and clarity. We regret that we cannot answer each one. Join the discussion of any article at www.SciAm.com/sciammag

SCIENTIFIC AMERICAN®

Established 1845

PRESIDENT: Steven Yee
MANAGING DIRECTOR, INTERNATIONAL: Kevin Hause
VICE PRESIDENT: Frances Newburg

MANAGING DIRECTOR, CONSUMER MARKETING: Christian Dorbandt
ASSOCIATE DIRECTOR, CONSUMER MARKETING: Anne Marie O'Keefe
SENIOR MARKETING MANAGER/RETENTION: Catherine Bussey
FULFILLMENT AND DISTRIBUTION MANAGER: Rosa Davis

VICE PRESIDENT AND PUBLISHER: Bruce Brandfon
DIRECTOR, GLOBAL MEDIA SOLUTIONS: Jeremy A. Abbate
SALES DEVELOPMENT MANAGER: David Tirpack
SALES REPRESENTATIVES: Gary Bronson, Jeffrey Crennan, Thomas Nolan, Stan Schmidt

PROMOTION MANAGER: Diane Schube
RESEARCH MANAGER: Aida Dadurian
PROMOTION DESIGN MANAGER: Nancy Mongelli

VICE PRESIDENT, FINANCE, AND GENERAL MANAGER: Michael Florek
BUSINESS MANAGER: Marie Maher
MANAGER, ADVERTISING ACCOUNTING AND COORDINATION: Constance Holmes

DIRECTOR, SPECIAL PROJECTS: Barth David Schwartz

DIRECTOR, ANCILLARY PRODUCTS: Diane McGarvey
PERMISSIONS MANAGER: Linda Hertz

How to Contact Us

SUBSCRIPTIONS

For new subscriptions, renewals, gifts, payments, and changes of address: U.S. and Canada, 800-333-1199; outside North America, 515-248-7684 or www.SciAm.com

REPRINTS

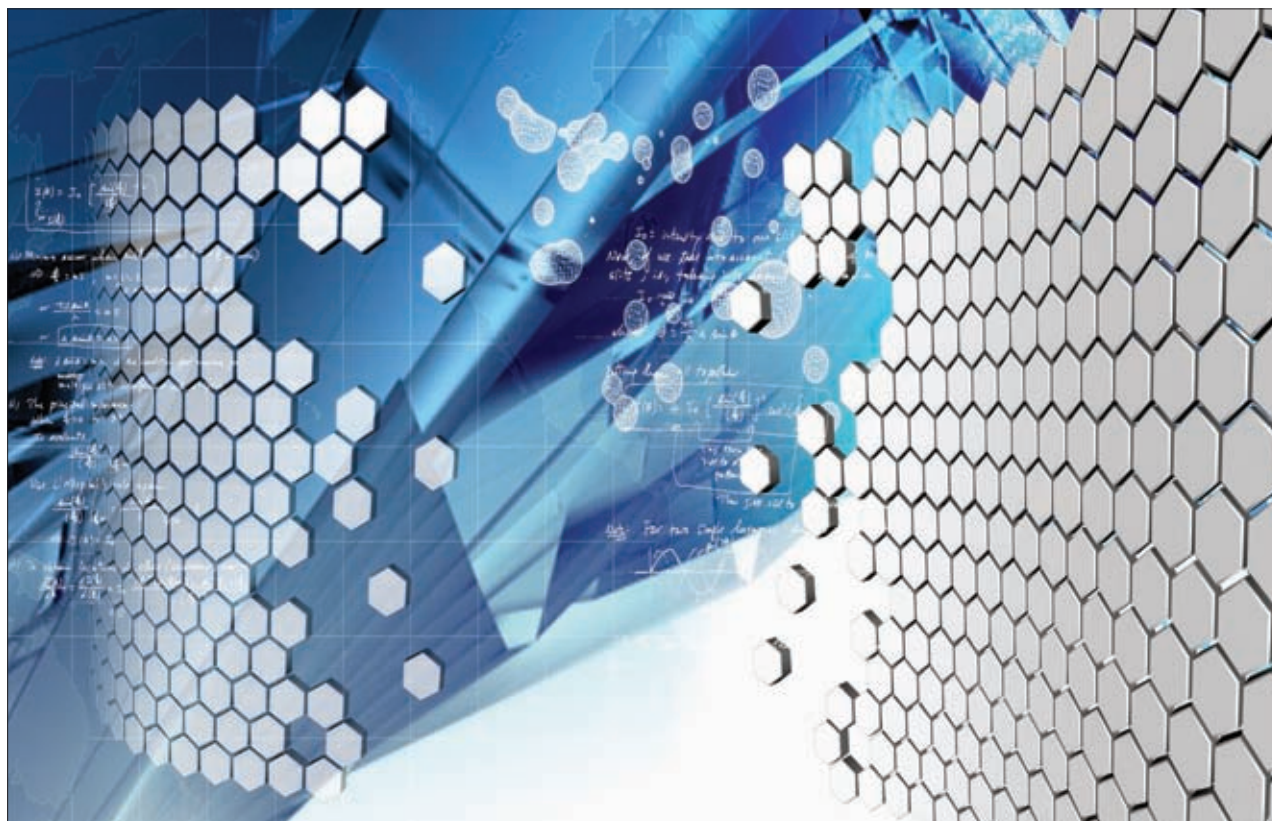
To order reprints of articles: Reprint Department, Scientific American, 415 Madison Ave., New York, NY 10017-1111; 212-451-8877, fax: 212-355-0408; reprints@SciAm.com

PERMISSIONS

For permission to copy or reuse material: Permissions Department, Scientific American, 415 Madison Ave., New York, NY 10017-1111; www.SciAm.com/permissions or 212-451-8546 for procedures. Please allow three to six weeks for processing.

ADVERTISING

www.SciAm.com has electronic contact information for sales representatives of Scientific American in all regions of the U.S. and in other countries.



Copyright © 2000-2008 Digitvision/Imagine.

Chaos Made Clear

It has been called the third great revolution of 20th-century physics, after relativity and quantum theory. But how can something called chaos theory help you understand an *orderly* world? What practical things might it be good for? What, in fact, is chaos theory? **Chaos** takes you to the heart of chaos theory as it is understood today. Your guide is Cornell University Professor Steven Strogatz, author of *Nonlinear Dynamics and Chaos*, the most widely used chaos theory textbook in America. In 24 thought-provoking lectures, he introduces you to a fascinating discipline that has more to do with your everyday life than you may realize.

Chaos theory—the science of systems whose behavior is sensitive to small changes in their initial conditions—affects nearly every field of human endeavor, from astronomy to business. Professor Strogatz shows you the importance of this field and how it has helped us solve life's mysteries. You learn how chaos theory was discovered and investigate ideas such as the butterfly effect, iterated maps, and fractals. You also discover practical applications of chaos in areas like encryption, medicine, and space mission design.

This course is one of The Great Courses®, a noncredit recorded college lecture series from The Teaching Company®. Award-winning professors of a wide array of subjects in the sciences and the liberal arts have made more than 250 college-level courses that are available now on our website.

Chaos

Taught by Professor Steven Strogatz, Cornell University

Lecture Titles

- | | |
|--|--|
| 1. The Chaos Revolution | 12. Experimental Tests of the New Theory |
| 2. The Clockwork Universe | 13. Fractals—The Geometry of Chaos |
| 3. From Clockwork to Chaos | 14. The Properties of Fractals |
| 4. Chaos Found and Lost Again | 15. A New Concept of Dimension |
| 5. The Return of Chaos | 16. Fractals Around Us |
| 6. Chaos as Disorder—The Butterfly Effect | 17. Fractals Inside Us |
| 7. Picturing Chaos as Order—Strange Attractors | 18. Fractal Art |
| 8. Animating Chaos as Order—Iterated Maps | 19. Embracing Chaos—From Tao to Space Travel |
| 9. How Systems Turn Chaotic | 20. Cloaking Messages with Chaos |
| 10. Displaying How Systems Turn Chaotic | 21. Chaos in Health and Disease |
| 11. Universal Features of the Route to Chaos | 22. Quantum Chaos |
| | 23. Synchronization |
| | 24. The Future of Science |

Order Today!
Offer expires Friday, March 20, 2009

Chaos
Course No. 1333
24 lectures (30 minutes/lecture)

DVDs ~~\$254.95~~ **NOW \$69.95**

+ \$10 Shipping, Processing, and Lifetime Satisfaction Guarantee

Priority Code: 32394

ACT NOW!

1-800-TEACH-12
www.TEACH12.com/1sa



THE TEACHING COMPANY®
The Joy of Lifelong Learning Every Day™
GREAT PROFESSORS, GREAT COURSES, GREAT VALUE
GUARANTEED.™

Pravda on Education ■ Teddy Roosevelt's Navy ■ Mass-Produced Ironwork

Compiled by Daniel C. Schlenoff

FEBRUARY 1959

SOVIET STUDENTS—"When Premier Khrushchev recently called upon Soviet educators to strengthen ties between the schools and 'life,' [Yakov B.] Zel'dovich and [Andrei] Sakharov wrote a long letter to *Pravda* on the training of scientists-to-be at the secondary-school level. Their thesis is that boys and girls with mathematical or scientific talent spend too many years in ordinary schools in view of the fact that mathematicians and theoretical physicists are often most productive in their early twenties. They recommend segregating such students at 14 or 15 in schools that emphasize mathematics, physics and chemistry, perhaps to the virtual exclusion of humanities."

WHIP IT GOOD—"Men created supersonic shock waves millennia before their projectiles and aircraft broke through the sound barrier. It seems that the crack of a whip occurs when its tip exceeds the speed of sound, and not when leather slaps against leather. This fact is revealed by an experimental and theoretical study of bullwhip dynamics made at the Naval Research Laboratory in Washington, D.C., with the cooperation of a team of theatrical whipcrackers. High-speed photographs made at 4,000 frames per second demonstrated that the tip moved at about 1,400 feet per second: some 25 per cent faster than sound. Shadow pictures clearly showed shock waves flowing from the tip."

FEBRUARY 1909

GREAT WHITE FLEET—"In view of the bitter criticism with which it was assailed, when the proposal to send a fleet of sixteen battleships from the Atlantic to the Pacific coast was first made public, the return of this same fleet to Hampton Roads after a 42,000-mile cruise around the world, with every ship in first-class shape and the *morale* of officers and men greatly improved, is a tribute to the far-sighted sagacity which

projected the voyage. The spectacle of this most imposing array of first-class fighting ships, steaming in perfect order and on scheduled time from port to port across all the seven seas, has had the effect of raising the prestige of our navy in every quarter of the world. If any American imagined that the rapidly-increasing power and wealth of this country was regarded with suspicion, distrust, or active envy, surely the wholehearted cordiality with which this concrete expression of our strength was everywhere received will effectually banish the idea from his mind."

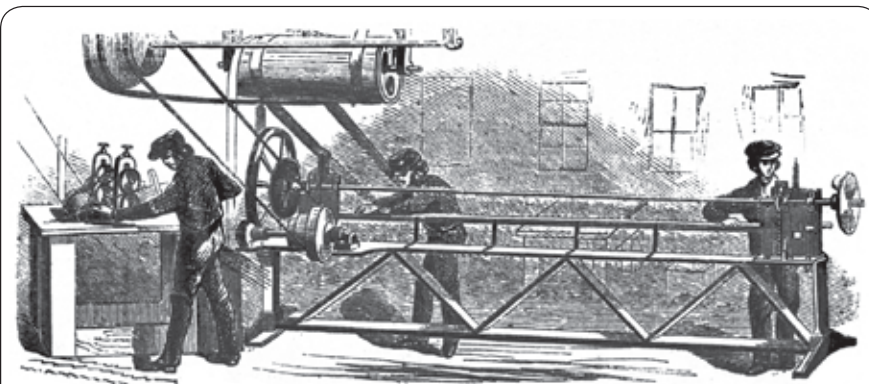
ASSMANN'S GASBAG—"At the present time, ascensions of the Assmann rubber balloons for meteorology to heights of 20,000 meters are not unusual. The most noteworthy improvement of the new method of sounding the air is the invention of Dr. Richard Assmann, director of the Royal Prussian Aeronautical Observatory. For the large balloons previously employed, some of which contained 500 cubic meters of gas, Dr. Assmann in 1901 substituted a much smaller one made of sheet rubber, which, when filled with hydrogen and sealed, rises until it is exploded by the internal expansion of the gas. The total weight of the 1,500-millimeter balloon, recording instruments, basket and cotton parachute is about 2,450 grammes."

FEBRUARY 1859

INDUSTRIAL AGE—"A comparatively new American art embraces very original manufactures of iron composite-work, such as railings, fences, household furniture, and such. It was known long ago that wrought iron was stronger and more flexible than any material employed in the arts, and that it was indestructible by the elements of the atmosphere, when protected with paint. But to forge it out of rods and bars into a great variety of forms, pleasing to the eye as well as useful and durable in character, was out of the question, owing to the great expense incurred for hand labor. The genius of the inventor was required to reduce wrought iron to practical purposes. The result is manufactures such as the New York Wire Railing Company [*see illustration*]."

➡ More images can be found at
www.SciAm.com/feb09

BEFORE ROOT BEER—"Dr. Bocker, of Bonn, on the Rhine, well known for his experiments on the digestion of articles of food, has discovered that sarsaparilla has none of those wondrous purifying properties usually attributed to it, and that it is a useless and expensive hospital drug. This but confirms the opinion which has been previously expressed through our columns."



INDUSTRIAL REVOLUTION and ironwork: The New York Wire Railing Company, 1859

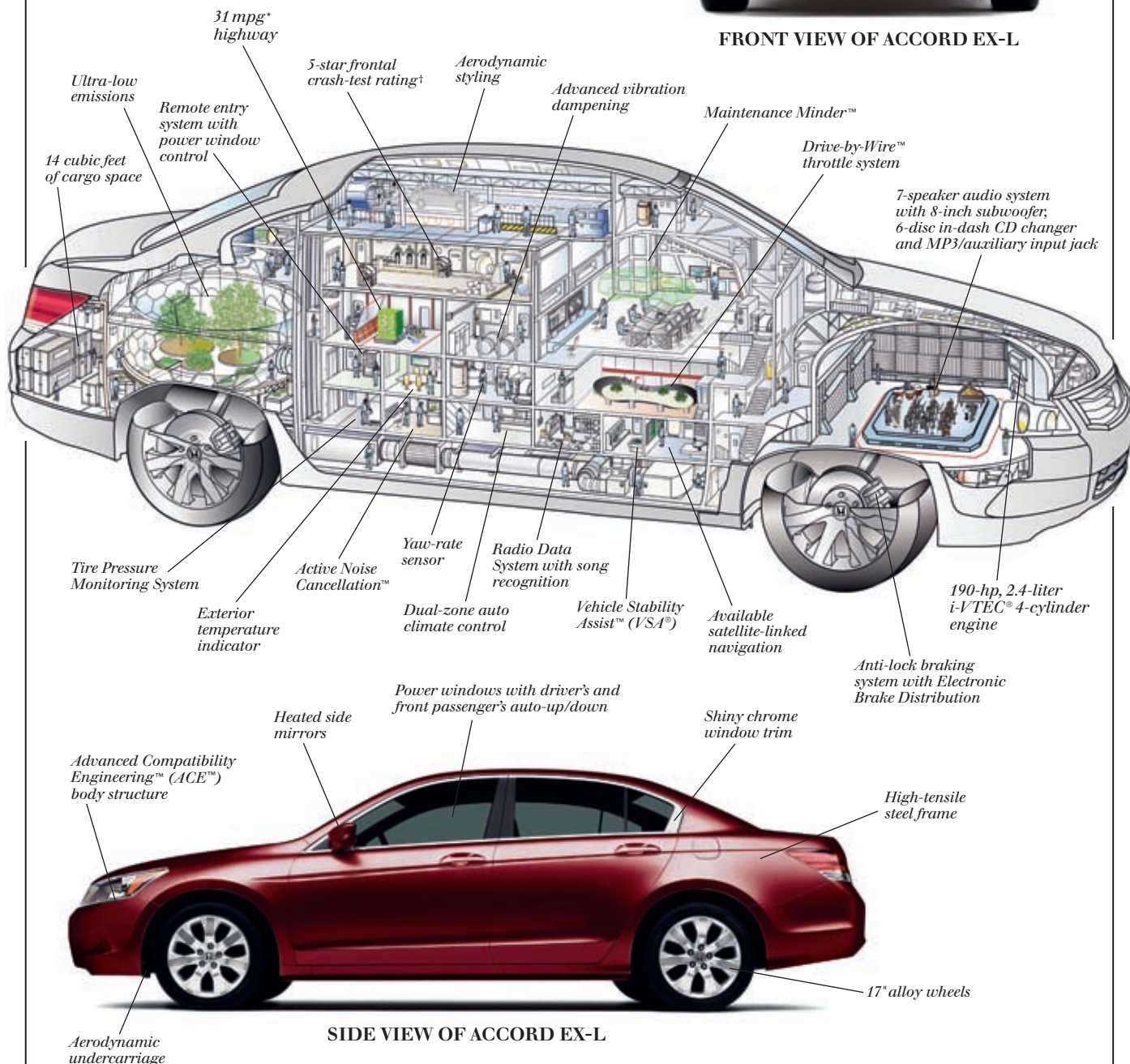
Efficient machines

LONG A SYMBOL OF EFFICIENCY AND QUALITY, over 10 million Accords have been sold in the United States since 1976. Building from a rich racing heritage and clever research, Honda engineers deftly combine efficiency, safety and performance in one place. The 2009 Accord.

Basque Red Pearl



FRONT VIEW OF ACCORD EX-L



SIDE VIEW OF ACCORD EX-L

It's all we know, all in one place. The 31-mpg* Accord.



*EPA-estimated hwy mpg based on 4-cylinder, 5-speed manual transmission. Use for comparison purposes only. Actual mileage will vary. EX-L Sedan model shown. †Based on 5-star frontal crash ratings. Government star ratings are part of the National Highway Traffic Safety Administration's (NHTSA's) New Car Assessment Program (www.safercar.gov). honda.com 1-800-33-Honda ©2008 American Honda Motor Co., Inc.

Wave Menace ■ Rabies Survival ■ Bioweapons Risk ■ Iberian Gene Mix

Edited by Philip Yam

■ No Relief from Tsunami Threat

In the devastating wake of the 2004 Indian Ocean tsunami, scientists rushed to investigate its cause and the potential for another killer wave [see “Tsunami: Wave of Change”; SciAm, January 2006]. They found that the tsunami resulted from a magnitude 9.2 earthquake off Sumatra’s western coast—specifically, at the Sunda megathrust, where one tectonic plate is diving below another. Scientists had conjectured two strong earthquakes there in 2007 might have relieved pent-up energy, thereby preventing another major quake.

Unfortunately, California Institute of Technology researchers and their colleagues have bad news. They analyzed satellite radar and GPS station data, along with coral reef growth patterns and other historical geologic records to assess how much of the area ruptured as compared with past activity. Evidently, the 2007 events released only a quarter of the stress trapped within. The teams report in the December 4 *Nature* and the December 12 *Science* that another tsunami-unleashing earthquake could occur there at any time. —Charles Q. Choi



TSUNAMI ruins in Banda Aceh, Indonesia.

■ Rabid Recoveries

In 2004 Jeanna Giese became the first person to survive a rabies infection without taking the vaccine. Rodney E. Willoughby, Jr., of the Medical College of Wisconsin saved her by inducing a coma and injecting her with antivirals [for his account, see “A Cure

for Rabies?”; SciAm, April 2007]. Last fall this “Milwaukee protocol” may have helped an eight-year-old Colombian girl and a teenage Brazilian boy beat the odds, too.

The girl began recovering while in her coma; before waking, however, she died of pneumonia, which her physicians say was unrelated to her rabies infection. The boy is recuperating, but whether the protocol worked is not completely clear: he had a partial course of the rabies treatment before showing symptoms. (Five others have recovered after such partial treatment.) Although the two cases may represent positive news for the Milwaukee protocol—research on it is controversial

and hard to conduct—Giese remains the only clear-cut success story for now.

■ Bioterror by 2013

Fear of deadly pathogens released as weapons of mass destruction has risen since the 9/11 attacks and the anthrax mailings [see “The Specter of Biological Weapons”; SciAm, December 1996, and “After the Anthrax”; SciAm, November 2008]. Such an event is more likely than a nuclear detonation, concludes a congressional commission, which in its December 2 report says a bioterror incident will likely occur somewhere by 2013. Cooperative efforts to secure pathogens and steer bioweapons scientists to peaceful activ-

ities has worked in the past, but such programs must be strengthened and extended, the commission states.

■ Religious Effects

The human genome can tell a story of ancestral migrations [see “Traces of a Distant Past”; SciAm, July 2008]. It also reveals the effects of religion and persecution. Investigators studied the genes of 1,140 males from around the Iberian Peninsula—specifically, their Y chromosome, which changes little from father to son.



EXPULSION and conversion of Jews show up in Iberian genes.

They found that 19.8 percent of the modern Iberian population has Sephardic Jewish ancestry. It probably reflects the 15th-century purge and conversion of Jews by Christians. Similarly, 10.6 percent have Moorish ancestry, probably a result of the Muslim conquest of the region in A.D. 711. The findings appear in the December 4, 2008, *American Journal of Human Genetics*.



BEATING RABIES: Jeanna Giese is wheeled out after her recovery.

DARIO MITTIERI/Getty Images (tsunami); GRANGER COLLECTION (19th-century engraving); MORRIS GASH/AP Photo (Giese and family)



He was a
hardworking farm boy.

She was an
Italian supermodel.

He knew he would
have just one chance
to impress her.

The fastest and easiest
way to learn ITALIAN.

Arabic • Chinese (Mandarin) • Danish • Dutch • English (American) • English (British) • French • German • Greek • Hebrew • Hindi • Indonesian • Italian • Irish • Japanese • Korean • Latin • Pashto • Persian (Farsi) • Polish • Portuguese (Brazil) • Russian • Spanish (Latin America) • Spanish (Spain) • Swahili • Swedish • Tagalog (Filipino) • Thai • Turkish • Vietnamese • Welsh

Rosetta Stone® brings you a complete language-learning solution, wherever you are: at home, in-the-car or on-the-go. You'll learn quickly and effectively, without translation or memorization. You'll discover our method, which keeps you excited to learn more and more.

- You'll experience **Dynamic Immersion**® as you match real-world images to words spoken by native speakers so you'll find yourself engaged and learn your second language like you learned your first.
- Our proprietary Speech Recognition Technology evaluates your speech and coaches you on more accurate pronunciation. You'll speak naturally.
- Only Rosetta Stone has **Adaptive Recall**™ that brings back material to help you where you need it most, for more effective progress.
- And Rosetta Stone includes **Audio Companion**™ so that you can take the Rosetta Stone experience anywhere you use a CD or MP3 player.

Innovative software. Immersive method. Complete mobility. It's the total solution. Get Rosetta Stone — **The Fastest Way to Learn a Language. Guaranteed.**®



SAVE 10%!

100% GUARANTEED
SIX-MONTH MONEY-BACK

Level 1	Reg. \$259	NOW \$233
Level 1&2	Reg. \$419	NOW \$377
Level 1,2&3	Reg. \$549	NOW \$494

©2008 Rosetta Stone Ltd. All rights reserved. Offer applies to Personal Edition only. Patent rights pending. Offer cannot be combined with any other offer. Prices subject to change without notice. Six-Month Money-Back Guarantee is limited to product purchases made directly from Rosetta Stone and does not include return shipping. Guarantee does not apply to an online subscription or to Audio Companion purchased separately from the CD-ROM product. All materials included with the product at the time of purchase must be returned together and undamaged to be eligible for any exchange or refund.

Call
(877) 214-6680

Online
RosettaStone.com/sas029

Use promotional code **sas029** when ordering.
Offer expires March 31, 2009.

RosettaStone® 

MEDICINE

Spreading the Health

Repositories for donated, unused drugs still face hurdles **BY JESSICA WAPNER**

Americans spend some \$200 billion annually on prescription drugs. Since 1997, in an effort to keep a lid on costs, 37 states have enacted legislation allowing patients, their families and health care facilities to recycle good, unused pills through local pharmacies for donation to patients lacking sufficient insurance. Thousands of patients could in principle benefit from these “drug repository” laws. But as well intentioned as these efforts are, practical problems have prevented widespread implementation of such programs.

The guidelines for these laws, which began thanks to the lobbying efforts of families of cancer patients, are fairly consistent throughout the country. Donated medications must be in sealed, tamper-evident packaging and usually must be within no more than six months of their expiration date. Pharmacies are not held liable should the drug’s next owner come to unexpected harm from the medication. Some repositories accept cancer drugs only; others take all prescriptions (minus narcotics and sleep aids). Some states accept unused pills from home medicine cabinets, whereas others, as a safety measure, permit donations only from professional facilities such as nursing homes.

Under the rules, Iowa collected more than 300,000 pills with a retail value of approximately \$290,000 in 2007 and distributed them to some 780 patients. Recycling medicines from Tulsa-area nursing

homes saves Oklahoma about \$120,000 a year. These successes, though, are small when compared with the potential of the practice. According to the American Cancer Society, as of June 2008 only about one third of the states with repository laws had up-and-running programs.

Part of the problem is money: pharmacies accepting donations do not want to incur the cost of hazardous waste disposal if the drugs go unused. With no reimbursement code for handling and processing donated meds, pharmacies have to be

established. “We don’t give any drug to anybody without knowing exactly where it’s been at all times,” says Roger Lyons, a private hematologist and oncologist in San Antonio, who regards the process as akin to filling prescriptions through the Internet or foreign pharmacies. “I am ultimately responsible for making sure a patient under my care gets the right medicine, so I’m not taking the risk.” Lyons also sees little need for repositories: “There are very few patients for whom we can’t get free drugs if they can’t otherwise afford it.”

The inability to ensure a ready supply is also problematic. Doug Englebert, who oversees Wisconsin’s drug repository program, notes that patients could suffer a potentially harmful gap in treatment if a pharmacy has a donated drug one month but not the next. Physicians, he says, “might have concerns with a repository because it’s not a guaranteed supply.”

Englebert cites some of the legal demands as hampering the usefulness of these programs. For example, the exclusion of drugs due to expire in less than

six months, which dramatically reduces the supply of eligible donations, may be overly cautious because many of the medicines would be claimed and used well within that time frame. Because the tamper-evident seal cannot be broken, even a nearly full bottle cannot be given. The requirement essentially limits donations to pills sealed in blister packs—otherwise known in the industry as unit-dose packaging.

“There are very few medications that



CHARITABLE PRESCRIPTIONS: State-legislated programs for the donation of unused drugs have seen limited success.

willing to operate as a repository on a completely charitable basis. And despite the letter of the law, many pharmacists fear lawsuits if the drugs prove faulty. Storing the drugs, especially when refrigeration is required, also poses its own issues and costs.

Physicians themselves have felt reluctant to steer patients toward the repositories. Many consider donated drugs too risky because their pedigree cannot be es-



CLIMATE CHANGE

Global Risks, Challenges & Decisions

COPENHAGEN 2009, 10-12 March

Experts and decision makers from more than 70 countries will gather to provide a scientific basis for urgent political action on climate change.

1.000+ international experts

57 multidisciplinary sessions

15 plenary speakers

1 synthesis report

handed over to the participants at the COP15

Read more and register at: www.climatecongress.ku.dk

"In my view the IARU International Scientific Congress on Climate Change will be a pivotal event that will contribute with a scientifically solid foundation to the political decisions to be made during the UN Climate Conference, COP15, in Copenhagen in December 2009"

Prime Minister of Denmark, Mr. Anders Fogh Rasmussen



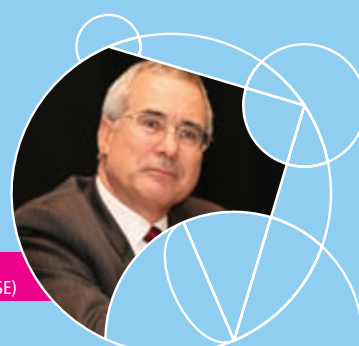
"The challenge before us is not only a large one, it is also one in which every year of delay implies a commitment to greater climate change in the future."

Dr. Rajendra Pachauri, Chairman of the IPCC



"We know climate change is the greatest market failure the world has ever seen. This Congress will be a significant contribution in articulating the latest evidence on how large the risks are and what the response should be."

Professor Lord Nicholas Stern,
The London School of Economics and Political Science (LSE)



The Star Sponsors of the IARU International Scientific Congress on Climate Change:



Principal Media Partner

TIME

Global Science Media Partner

SCIENTIFIC
AMERICAN

Media Partner

NATIONAL
GEOGRAPHIC

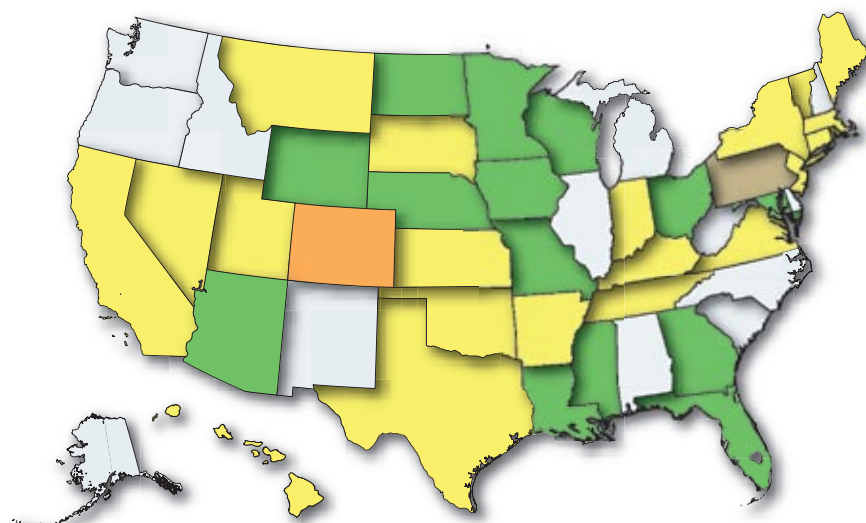
The University of Copenhagen is organising the scientific congress in cooperation with the partners in the International Alliance of Research Universities (IARU):



INTERNATIONAL ALLIANCE OF
RESEARCH UNIVERSITIES

Australian National University, ETH Zürich, National University of Singapore, Peking University, University of California - Berkeley, University of Cambridge, University of Copenhagen, University of Oxford, The University of Tokyo, Yale University

STATES THAT PERMIT DRUG DONATIONS



WHO CAN DONATE

- Any person, drug manufacturer or health care institution
- Professional facilities and businesses only, such as nursing homes, pharmacies and hospitals
- Patients or their families
- Health care institution
- No donation programs

Colorado, Florida, Minnesota, Nebraska and Pennsylvania permit cancer-related donation programs only; Wisconsin's is for chronic diseases only.

are in unit-dose packaging,” Englebert remarks, “and so therefore very few that are eligible for donation.” In addition, the lack of funding renders many programs cum-

bersome. Without databases of participating pharmacies and their current inventory, for instance, would-be recipients need to call every registered outlet to inquire

whether their prescription is available.

To increase the utility of the repository laws, health counselors, pharmacists and volunteers have deployed various strategies. Some clinics are incorporating repositories into their ongoing patient assistance programs. Other efforts focus on specific medications, such as high-cost cancer drugs to which patients often prove intolerant.

Education is also key: pharmacy counters could provide information about what consumers can do with unused medications. And as Englebert points out, tackling packaging issues up front—such as increased use of blister seals—might help satisfy security requirements.

Many experts and patient advocates remain optimistic about drug repositories. Sarah Barber, who is a senior policy analyst at the American Cancer Society, notes that the nationwide trend indicates a definite need. These programs, she thinks, “will become much easier and much more usable in the future.”

Jessica Wapner, based in New York City, writes frequently about health care issues.

JOHNNY JOHNSON, SOURCE: NATIONAL COUNCIL OF STATE LEGISLATURES

OPTIMIZATION

Detours by Design

How closing streets and removing traffic lights speed up urban travel **BY LINDA BAKER**

Conventional traffic engineering assumes that given no increase in vehicles, more roads mean less congestion. So when planners in Seoul tore down a six-lane highway a few years ago and replaced it with a five-mile-long park, many transportation professionals were surprised to learn that the city's traffic flow had actually improved, instead of worsening. “People were freaking out,” recalls Anna Nagurney, a researcher at the University of Massachusetts Amherst, who studies computer and transportation networks. “It was like an inverse of Braess's paradox.”

The brainchild of mathematician Dietrich Braess of Ruhr University Bo-

chum in Germany, the eponymous paradox unfolds as an abstraction: it states that in a network in which all the moving entities rationally seek the most efficient route, adding extra capacity can actually reduce the network's overall efficiency. The Seoul project inverts this dynamic: closing a highway—that is, reducing network capacity—improves the system's effectiveness.

Although Braess's paradox was first identified in the 1960s and is rooted in 1920s economic theory, the concept never gained traction in the automobile-oriented U.S. But in the 21st century, economic and environmental problems are bringing new

scrutiny to the idea that limiting spaces for cars may move more people more efficiently. A key to this counterintuitive approach to traffic design lies in manipulating the inherent self-interest of all drivers.

A case in point is “The Price of Anarchy in Transportation Networks,” published last September in *Physical Review Letters* by Michael Gastner, a computer scientist at the Santa Fe Institute, and his colleagues. Using hypothetical and real-world road networks, they explain that drivers seeking the shortest route to a given destination eventually reach what is known as the Nash equilibrium, in which no single driver can do any better by

changing his or her strategy unilaterally. The problem is that the Nash equilibrium is less efficient than the equilibrium reached when drivers act unselfishly—that is, when they coordinate their movements to benefit the entire group.

The “price of anarchy” is a measure of the inefficiency caused by selfish drivers. Analyzing a commute from Harvard Square to Boston Common, the researchers found that the price can be high—selfish drivers typically waste 30 percent more time than they would under “socially optimal” conditions.

The solution hinges on Braess’s paradox, Gastner says. “Because selfish drivers optimize a wrong function, they can be led to a better solution if you remove some of the network links,” he explains. Why? In part because closing roads makes it more difficult for individual drivers to choose the best (and most selfish) route. In the Boston example, Gast-

ner’s team found that six possible road closures, including parts of Charles and Main streets, would reduce the delay under the selfish-driving scenario. (The street closures would not slow drivers if they were behaving unselfishly.)

Another kind of anarchy could actually speed travel as well—namely, a counterintuitive traffic design strategy known as

shared streets. The practice encourages driver anarchy by removing traffic lights, street markings, and boundaries between the street and sidewalk. Studies conducted in northern Europe, where shared streets are common, point to improved safety and traffic flow.

The idea is that the absence of traffic regulation forces drivers to take more responsibility for their actions. “The more uncomfortable the driver feels, the more he is forced to make eye contact on the street with pedestrians, other drivers and to intuitively go slower,” explains Chris Conway, a city engineer with Montgomery, Ala. Last April the city converted a signalized downtown intersection into a European-style cobblestone plaza shared by cars, bikes and pedestrians—one of a handful of such projects that are springing up around the country.

Although encouraging vehicular



MEAN STREETS: Urban travel is slow and inefficient, in part because drivers act in self-interested ways.

MAREMAGNUM Getty Images

A USER'S GUIDE FOR THE BRAIN



**WRAP YOUR
MIND
AROUND THIS...**

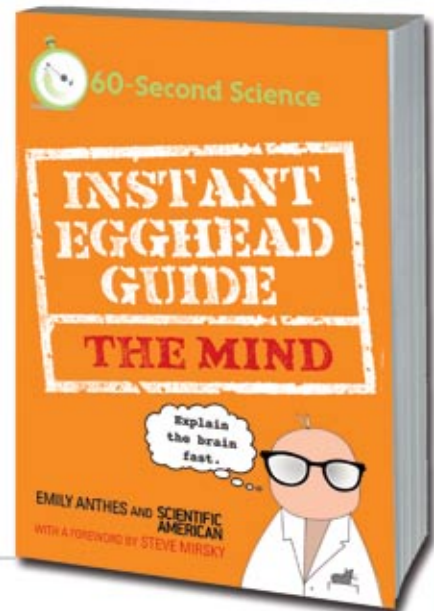
DECISION MAKING

Recent research on decision-making focuses on serotonin. In one study, researchers determined that those with reduced serotonin levels were less likely to cooperate or share money fairly. **INCREASED SEROTONIN WILL LEAD YOU TO BUY THIS BOOK.**



60-Second Science & SCIENTIFIC AMERICAN bring you **INSTANT EGGHEAD GUIDES: We explain BIG ideas, fast.**

Free excerpts at: www.instantegghead.com



chaos seems at odds with the ideas presented in the price of anarchy study, both strategies downplay the role of the individual driver in favor of improved outcomes for everyone. They also suggest a larger transportation niche for bicycles and pedestrians. As the Obama administration prepares to invest in the biggest public works project since the construction of the interstate highway system, the notion that fewer, more inclusive roads yield better results is especially timely.

Linda Baker is based in Portland, Ore.

Faster Streets with Less Parking

New strategies in parking management could also improve urban traffic flow, remarks Patrick Siegman, a principal with Nelson/Nygaard Consulting Associates in San Francisco, a transportation-planning firm. In a misguided effort to reduce congestion, planners in the 1950s required developers to provide a minimum number of free parking spaces—a strategy that “completely ignored” basic economics, Siegman says, referring to how lower prices increase demand.

Now limited urban space and concerns about global warming are inspiring city planners to eliminate these requirements. In San Francisco, for example, developers must restrict parking to a maximum of 7 percent of a building’s square footage, a negligible amount. Although downtown employment has increased, traffic congestion is actually declining, Siegman says. With fewer free spaces to park, drivers seem to be switching modes, relying more on mass transit, cycling and just plain walking.

ECOLOGY

Plague in the Prairie

To contain U.S. outbreaks, it's Kazakhstan's giant gerbils to the rescue **BY PAUL VOSEN**

Plague conjures images of Gothic horror—rough wooden carts piled high with pestilent bodies—but it is more than a medieval memory. The disease, caused by the bacterium *Yersinia pestis*, kills several hundred people every year by attacking the lungs, lymph nodes or blood. Less obviously, plague also ravages wildlife around the world.

Introduced to the U.S. a century ago, it is creeping into the upper Midwest, wiping out prairie dogs and threatening the black-footed ferret, one of North America’s rarest species. Confined to rural regions, the disease so far is not a major threat to people—only a few Americans die from it annually. But things could change if the bacterium spreads to urban-loving rodents such as rats. Now some researchers think that another species could provide the information needed to contain plague’s spread in the U.S.: the giant gerbils of Kazakhstan.

Inhabitants of the vast steppes of Central Asia, the gerbils grow to one foot in length. They are natural hosts for *Yersinia*, and many researchers believe that the plague bacterium, carried by fleas hitching rides on the Mongols centuries ago, spread from these gerbils. Until World

War II, plague killed scores of people every year in Kazakhstan. “Whole villages were being wiped out,” recounts Stephen Davis, an Australian researcher who recently joined Yale University’s School of Public Health.

The former Soviet Union, which controlled the region at the time, cracked down on the disease: beginning in 1949, it sent teams into the steppe to collect gerbils and fleas to determine the extent of the outbreaks. They fumigated infected burrows with insecticide, killing the fleas but sparing the gerbils. Plague fatalities among humans dropped to a few cases a year. (Antibiotics can cure infected individuals

if they are treated promptly.) The control program has continued relatively unchanged to this day, although funding from the Kazakh government has tailed off in recent years.

All these data have become a rich resource, says Michael Begon, an ecologist at the University of Liverpool in England. The archive, housed by the Kazakh Scientific Center for Quarantine and Zoonotic Diseases in Almaty, came to the attention of Western researchers in 1996, when Herwig Leirs, an ecologist at the University of Antwerp in Belgium, reviewed a funding application from the center. Leirs was astonished. “We said, ‘This is potentially a gold mine for doing research,’” he recalls. The archive was “literally in big books, handwritten in thick ledgers,” stored in Almaty and 12 regional stations, he describes. Using a small set of plague data, Davis, Begon, Leirs and others published their initial findings in *Science* in 2004, namely, that when the gerbil population exceeded a certain threshold, plague outbreaks occur two years later.

Scientists are now focused on developing an early-warning system using this threshold, Begon says. This group includes Davis, who published a paper in *Nature*



PLAGUE VICTIM: The endangered black-footed ferret is getting infected via prairie dogs.

JIM BRANDENBURG/Minden Pictures

As it turns out, ideas are
our planet's most precious resource.

Submit your ideas to make the world a better place and you could be
the winner of The "Why Not?" Innovation Experience.

toyota.com/whynotinnovate



TOYOTA

NO PURCHASE NECESSARY TO ENTER OR WIN. Open only to licensed drivers who are legal residents of the 50 United States and D.C., 21 years of age or older. Entries must be received by 3/31/09. Subject to additional terms and conditions. See official rules for details, available at www.toyota.com/whynotinnovate ©2009

last summer that furthered the theoretical understanding of how plague spreads by employing percolation theory. In physics, it can explain, for example, how a fluid spreads through a porous medium—without enough connecting pores, the fluid remains in pockets. In epidemiology, percolation theory can model how disease spreads in situations that do not have random mixing, such as the widely spaced burrows of the gerbils.

The fixed nature of the burrows means that not every infected site may have to be fumigated. Such a targeted approach could require less funding—perfect for Central Asia and potentially useful information about the spread in other rodents, such as prairie dogs.

The U.S. has had little endemic pockets of plague since 1898, when it arrived from Asia. In the past two years, plague outbreaks have pushed into South Dakota, says Christopher Brand of the National Wildlife Health Center, who coordinated a November symposium on the issue. Having no natural defense, prairie dogs are especially vulnerable to plague; the mortality rate hovers around 90 percent. Already the disease has killed one third of the prairie dog population in South Dakota's Conata Basin.

Conservationists are in particular worried about plague's effect on the black-footed ferret, an endangered species that preys on prairie dogs, its primary food source. Desperate to save North America's

only native ferret, the U.S. Fish and Wildlife Service is spraying prairie dog burrows with insecticide. The service has also begun a capture-and-release vaccination program for the ferrets.

These labor-intensive tasks could cost less if the threshold model derived from the giant gerbils holds up; Kazakh scientists are now testing the concept. Davis says he is "quite encouraged" by the similarity in data between the prairie dogs and the gerbils. To further compare notes, he and other concerned researchers plan to meet in Kazakhstan this spring, as the harsh winter thaws and the gerbils emerge from their burrows.

Paul Voosen is based in New York City.

PHYSICS

Quantum Afterlife

A way for quantum benefits to survive after entanglement ends **BY CHARLES Q. CHOI**

"Sooky action at a distance" is how Albert Einstein famously derided the concept of quantum entanglement—where objects can become linked and instantaneously influence one another regardless of distance. Now researchers suggest that this spooky action in a way might work even beyond the grave, with its effects felt after the link between objects is broken.

In experiments with quantum entanglement, which is an essential basis for quantum computing and cryptography, physicists rely on pairs of photons. Measuring one of an entangled pair immediately affects its counterpart, no matter how far apart they are theoretically. The current record distance is 144 kilometers, from La Palma to Tenerife in the Canary Islands.

In practice, entanglement is an extremely delicate condition. Background disturbances readily destroy the state—a bane for quantum computing in particular, because calculations are done only as long as the entanglement lasts. But for the first time, quantum physicist Seth Lloyd of

the Massachusetts Institute of Technology suggests that memories of entanglement can survive its destruction. He compares the effect to Emily Brontë's novel *Wuthering Heights*: "the spectral Catherine communicates with her quantum Heathcliff as a flash of light from beyond the grave."

The insight came when Lloyd investigated what happened if entangled photons were used for illumination. One might suppose they could help take better pictures. For instance, flash photography shines light out and creates images from photons that are reflected back from the object to be imaged, but stray photons from other objects could get mistaken for the returning signals, fuzzing up snapshots. If the flash emitted entangled photons instead, it would presumably be easier to filter out noise signals by matching up returning photons to linked counterparts kept as references.

Still, given how fragile entanglement is, Lloyd did not expect quantum illumination to ever work. But "I was desperate," he recalls, keen on winning funding from

a Defense Advanced Research Projects Agency's sensor program for imaging in noisy environments. Surprisingly, when Lloyd calculated how well quantum illumination might perform, it apparently not only worked, but "to gain the full enhancement of quantum illumination, all entanglement must be destroyed," he explains.

Lloyd admits this finding is baffling—and not just to him. Prem Kumar, a quantum physicist at Northwestern University, was skeptical of any benefits from quantum illumination until he saw Lloyd's math. "Everyone's trying to get their heads around this. It's posing more questions than answers," Kumar states. "If entanglement does not survive, but you can seem to accrue benefits from it, it may now be up to theorists to see if entanglement is playing a role in these advantages or if there is some other factor involved."

As a possible explanation, Lloyd suggests that although entanglement between the photons might technically be completely lost, some hint of it may remain in-

tact after a measurement. "You can think of photons as a mixture of states. While most of these states are no longer entangled, one or a few remain entangled, and it is this little bit in the mixture that is responsible for this effect," he remarks.

If quantum illumination works, Lloyd suggests it could boost the sensitivity of radar and x-ray systems as well as optical telecommunications and microscopy by a millionfold or more. It could also lead to stealthier military scanners because they could work even when using weaker signals, making them easier to conceal from adversaries. Lloyd and his colleagues detailed a proposal for practical implementation of quantum illumination in a paper submitted in 2008 to *Physical Review Letters* building off theoretical work presented in the September 12 *Science*.

Actually proving this effect

may be the real challenge. The easy part is creating entangled photons: just shoot light through a special, "downconverting" crystal that acts as a beam splitter; it produces separate yet linked rays. One ray illuminates the object, and the other serves as a reference. The returning and reference beams then are merged together (basically,

by making them go through a splitter in reverse); the photons that were entangled should be more likely to recombine, or "upconvert." But any experiment to prove that quantum illumination can boost the sensitivity of imaging has to use weak signals, and creating materials capable of up-converting faint beams with high efficiency is technically daunting, Kumar says. Still, Lloyd predicts experimental tests of this scheme might come later this year.

Besides boosting imaging sensitivity, the effect might confer benefits on quantum computing or quantum cryptography, Kumar suspects. "The quantum world is quite exotic and complex, and this shows there are surprises there that lurk around corners all the time," he says.

Charles Q. Choi is a frequent contributor.



LIGHTS, CAMERA ... ENTANGLEMENT! In theory, images could be dramatically sharper if flash units used entangled photons.

MIKE POWELL/Getty Images

HERE'S TO THE VISIONARIES WHO COME PREPARED FOR ANYTHING. PARTICULARLY ANYTHING TEARING, BREAKING OR LEAKING.

GORILLA TAPE
For The Toughest Jobs On Planet Earth.
INCREDIBLY STRONG

FOR THE TOUGHEST JOBS ON PLANET EARTH.® 1-800-966-3458 gorillatough.com Made in U.S.A

© 2009 Gorilla Glue Company 09_T1HD

MATHTUTORDVD.COM
Press Play For Success

Introducing the Ultimate Physics 2 Tutor Volume 1

Ultimate Physics 2 Tutor Volume 1
18 Hour Video Course
3 DVD Set
A+
CAUTION!
THE USE OF THIS DVD WILL CAUSE INCREASED UNDERSTANDING OF PHYSICS 2 AND WILL LEAD TO HIGHER GRADES

Only \$39.99
18 Hour DVD Course

DVD Includes:

- Thermometers/Temperature Scales
- Expansion & Contraction of Solids/Liquids
- Kinetic Theory of Gases
- Heat/Latent Heat/Phase Change
- Heat Transfer by Convection/Radiation/Conduction
- Work/Heat/PV Diagrams
- The First Law of Thermodynamics
- Heat Engines/the Second Law of Thermodynamics
- Refrigerators
- Entropy

Also Available: Courses in Basic Math, Algebra, Trig, Calculus 1-3, Physics 1, Probability & more!

#1 Rated Math & Physics Tutorial DVDs
Visit MathTutorDVD.com to view sample videos

877-MATH-DVD
Visit: MathTutorDVD.com/sciam

PALEONTOLOGY

Shell Game

Turtle origins come out of hiding **BY KATE WONG**

Vertebrate animals come in all shapes and sizes. But some have evolved truly bizarre forms. With beaks instead of teeth and shells formed by the ribs and other bits, turtles surely rank among the strangest of our backboned brethren. Indeed, paleontologists have long puzzled over how turtles acquired their odd traits and who their closest relatives are.

Previously, much of what researchers knew about turtle origins derived from fossils of *Proganochelys* from Germany. Based on that creature, with its heavily built shell and spiked plates covering the

illuminates how their trademark armature took shape. Dating back to 220 million years ago, this transitional creature, named *Odontochelys semitestacea* (“half-shelled turtle with teeth”), is the oldest and most primitive turtle on record. Researchers led by Chun Li of the Chinese Academy of Sciences in Beijing describe the fossil in the November 27, 2008, issue of *Nature*.

Odontochelys possesses a plastron—the flat, lower half of the shell that protects the animal’s soft belly—but lacks the domed upper half. What this suggests, Li and his colleagues say, is that the shell evolved from

with the ribs and backbones to form a carapace. In fact, last October researchers writing in the *Proceedings of the Royal Society B* reported on a 210-million-year-old turtle fossil from New Mexico believed to support that hypothesis.

But critics have countered that findings from turtle embryology hinted that the backbones and ribs alone morphed to make a shell. *Odontochelys* bolsters the theory that ribs flattened and spread to form the top of the shell.

The absence of osteoderms in *Odontochelys* also challenges the idea that turtles are closely related to pareiasaurs. Taken together with molecular data, the new evidence aligns the shelled vertebrates with another group of reptiles, the diapsids.

Some aspects of the discovery team’s interpretation of *Odontochelys* have alternative explanations, however. In a commentary accompanying the *Nature* paper, paleontologists Robert Reisz and Jason Head of the University of Toronto Mississauga argue that the animal did have an upper shell, just one that had not fully ossified. If correct, their supposition would suggest that the form of this animal’s shell, rather than being a primitive intermediate, is a specialized adaptation. It turns out that aquatic turtles often have smaller, more delicate upper shells compared with their landlubber counterparts, as seen in sea turtles and snapping turtles.

Thus, rather than showing that turtles evolved in the water, Reisz and Head contend, *Odontochelys* may signal an early invasion of the water by turtles that originated on terra firma. “The morphology of *Odontochelys* suggests that this story is more complex and more interesting than suggested” by Li and his co-authors, Reisz remarks. “We feel that *Odontochelys* is not the final answer; it is instead one more piece in the fascinating puzzle of turtle origins.”



TRANSITIONAL TURTLE *Odontochelys semitestacea*, the oldest turtle fossil yet, has a fully formed lower shell, or plastron, but lacks a fully formed upper shell.

neck and tail, researchers had proposed that turtles were kissing cousins of a group of extinct armored reptiles known as pareiasaurs. They also suggested that the first turtles lived on land, where a shield is a useful defense for a slow-footed creature. *Proganochelys* furnished no clues to how the turtle shell evolved, however, because its own carapace is fully formed.

A newfound fossil from southwestern China’s Guizhou Province paints a rather different picture of the origin of turtles and

the bottom up. In addition, the deposits that yielded the fossil indicate that this animal lived in a marine environment. If so, the plastron would have shielded the turtle’s underside from predators approaching from below.

Odontochelys also lacks osteoderms, bony plates in the skin that form the armor of reptiles such as crocodiles and dinosaurs. Some specialists had proposed that the turtle’s shell began as rows of osteoderms that, over millions of years, fused

NEUROBIOLOGY

Childhood Recovered

"Lazy eye" studies show how adult brains can be rewired back to youth **BY GARY STIX**

The pirate look is a time-honored way to fix children's "lazy eye": the patch over the good eye forces the weak one to work, thereby preventing its deterioration. Playing video games helps, too. The neural cells corresponding to both eyes then learn to fire in synchrony so that the brain wires itself for the stereo vision required for depth perception. Left untreated past a critical age, lazy eye, or amblyopia, can result in permanently impaired vision. New studies are now showing that this condition, which affects up to 5 percent of the population, could be repaired even past the critical phase.

What is more, amblyopia may provide insights into brain plasticity that could help treat a variety of other disorders related to faulty wiring, including schizophrenia, epilepsy, autism, anxiety and addiction. These ailments "are not neurodegenerative diseases that destroy part of the neural circuitry," notes Takao Hensch, a Harvard Medical School researcher. So if the defective circuits "could be stimulated in the right way, the brain could develop normally."

The recent findings have their roots in work from 10 years ago. Then, Hensch led a team that discovered the specific visual circuitry that induces a "critical period" during early life in which the two eyes must work together to establish the connections in the cortex underlying proper visual acuity. So-called parvalbumin basket cells release the neurotransmitter GABA, which puts the brakes on cell activity. But GABA and compounds that behave like it—the drug Valium, for one—can also trigger the critical phase. It is paradoxical that neurochemicals that turn cells off play a role in initiating a key developmental stage.

Hensch's discovery, along with the recognition of the important part played by the proteins and sugars that form a

matrix surrounding parvalbumin cells, has resulted in a set of recent experiments that demonstrate ways to reinstate the critical period in adult animals—and perhaps to map a path toward treatments. In 2006 a group led by Lamberto Maffei, a neurobiologist at the University of Pisa in Italy, injected an enzyme called chondroitinase into the visual cortex of adult rats with amblyopia to dissolve the extracellular matrix and restore the critical period. After patching a rat's good eye, the researchers witnessed the recovery of normal vision: cortical circuitry for both the left and right eyes were nudged into firing together, just as they are during the early phase of childhood development.

More recently, Hensch's team reported in *Cell* last summer on a protein that has the same effect as Valium in the developing visual cortex. Called Otx2, it has a role in the embryonic development of the head and becomes prominent again after birth, serving as the starting gun for the critical period. The protein actually travels from the retina to the visual cortex at the rear of the brain, perhaps because the visual cortex needs to wait for a signal from the eyes that it is ready to undergo maturation.

Hensch also presented work at the Society for Neuroscience annual meeting in November on adult mice with amblyopia that were genetically engineered to lack a recep-

SCIENCE FACULTY NYU Abu Dhabi

New York University (NYU) is establishing a new comprehensive liberal arts campus in Abu Dhabi, the capital of the United Arab Emirates. New York University Abu Dhabi (NYUAD), a partner campus of New York University New York (NYUNY), will consist of a highly selective liberal arts college (Arts, Humanities, Social Sciences, Sciences and Engineering), distinctive graduate programs, and a world-class Institute for advanced research, scholarship, and creative work. NYUNY and NYUAD will be integrally connected, together forming the foundation of a unique global network university, actively linked as well to NYU's study and research sites on five continents.

The Division of Science, Engineering, Technology and Mathematics at NYUAD is now recruiting faculty of exceptional quality in teaching and in professional accomplishment. The Division is specifically looking for professors of Biology, Chemistry, Computer Science, Neuroscience, Mathematics and Physics. Recruited faculty will start by teaching an innovative, three-semester foundational core, called the Science Foundations series. The Science Foundations series is especially designed to integrate basic concepts from mathematics, physics, chemistry, biology, computer science, and neuroscience and is required for all science majors at NYUAD. Science instruction at NYUAD will start in AY2010; however, applicants could start at NYUNY in September, 2009. Modern science laboratories will be constructed and become available for faculty research within a start-up phase spread over several subsequent years. Faculty may spend time at NYU in New York and at its other global campuses. The terms of employment are competitive compared to U.S. benchmarks and include housing and educational subsidies for children.

The review of applications will begin in January 2009 and will continue until the positions are filled. To apply for this position, please send ONE document only (pdf or word) via email to: nyuad.science@nyu.edu. This document should contain a cover letter (please address the letter to NYUAD Science Search Committee), a curriculum vitae, statements of teaching experience and research interests, contact information for references and representative publications. Electronic submissions are preferred, but you may also send a hard copy to: NYUAD Science Search Committee, New York University, 70 Washington Square South, Rm. 1242, New York, NY 10012. Information concerning the faculty, programs and facilities of NYU Abu Dhabi, can be obtained at: <http://nyuad.nyu.edu>



**NEW YORK UNIVERSITY
ABU DHABI**

NYU Abu Dhabi is an Equal Opportunity/Affirmative Action Employer.

SUBSCRIBE, RENEW OR
GIVE A GIFT ONLINE!



To give a gift subscription of
Scientific American:

www.SciAm.com/gift

To renew your subscription to
Scientific American:

www.SciAm.com/renew

To subscribe to
Scientific American Mind:

www.SciAmMind.com

To subscribe to
Scientific American Digital:

www.SciAmDigital.com

SCIENTIFIC AMERICAN
MIND SCIENTIFIC
AMERICAN

For additional subscription information, please
e-mail customer services at: sacust@sciam.com.
For information about our privacy practices,
please read our Privacy Notice at:
<http://www.sciam.com/privacy>.

NEWS SCAN

OPTICAL WORKOUT: Patching a child's good eye during a "critical period" may rewire the brain to restore visual acuity in the impaired eye.



tor on neurons for Nogo, a growth-inhibiting protein that originates in the myelin insulation around the neural wires called axons. In the experiment, suturing shut one of the two healthy eyes during the critical period induced amblyopia and its attendant decrease in visual acuity. When the sutures were removed, however, the mice that did not have the molecular brake of the Nogo receptor spontaneously regained their vision.

"This work is inspirational for me," remarks Dennis Levi, a neuroscientist at the University of California, Berkeley. "The future will be some kind of molecular intervention for amblyopia."

Such a future may not be far off. In fact, oral compounds may already exist on the pharmaceutical shelves. Last year Maffei's group found that the antidepressant Prozac can restore plasticity in the adult visual system of rats.

The ability to revert neural cells back

to their younger, plastic state could potentially be a treatment breakthrough. But fully restoring the brain's original sponge-like quality would nonetheless give clinicians pause. Turning the brain into malleable mush at the age of 30 would not be the best solution—some scientists think that an excess of plasticity, in fact, may lie at the root of conditions such as schizophrenia.

Some investigators are already exploring how far they can bring back plasticity and mend patients through environmental cues alone. In his own work, Levi found that after thousands of sessions in video game–like exercises—"kilo trials" as he calls these mini clinical trials—adult amblyopia patients achieved substantial improvements in visual acuity. Levi is already doing research with actual video games. Grand Theft Auto IV or Medal of Honor may retrain the brain in ways its developers never imagined.

Perils of a Badly Wired Brain

The research showing that neural systems can be forced back to an earlier, more pliable state may extend beyond treatments for "lazy eye." Schizophrenia may emerge from faulty signals transmitted during the critical developmental period, causing an excess of plasticity throughout life. Autistic children may suffer a surfeit of overexcited connections, another offshoot of errors in wiring that occur during this early-childhood window. Biochemicals similar to those in the visual system may be activated by auditory, olfactory and tactile signals. Adjusting their levels up or down in the central nervous system could conceivably treat a variety of disorders.

ERIN PATRICE O'BRIEN/Getty Images

OCEANS

Acid Bath

Carbon dioxide may be acidifying seawater faster than thought

BY CHARLES Q. CHOI

A lesser-known consequence of having a lot of carbon dioxide (CO₂) in the air is the acidification of water. Oceans naturally absorb the greenhouse gas; in fact, they take in roughly one third of the carbon dioxide released into the atmosphere by human activities. When CO₂ dissolves in water, it forms carbonic acid, the same substance found in carbonated beverages. New research now suggests that seawater might be growing acidic more quickly than climate change models have predicted.

Marine ecologist J. Timothy Wootton of the University of Chicago and his colleagues spent eight years compiling measurements of acidity, salinity, temperature and other data from Tatoosh Island off the northwestern tip of Washington State. They found that the average acidity rose more than 10 times faster than predicted by climate simulations.

Highly acidic water can wreak havoc on marine life. For instance, it can dissolve the calcium carbonate in seashells and coral reefs [see “The Dangers of Ocean Acidification,” by Scott C. Doney; *SCIENTIFIC AMERICAN*, March 2006]. In their study, published in the December 2 *Proceedings of the National Academy of*

Sciences USA, Wootton and his team discovered that the balance of ecosystems shifted: populations of large-shelled animals such as mussels and stalked barnacles dropped, whereas smaller-shelled species and noncalcareous algae (species that lack calcium-based skeletons) became more abundant. “I see it as a harbinger of the trends we might expect to occur in the future,” says oceanographer Scott C. Doney of the Woods Hole Oceanographic Institution, who did not participate in this study.

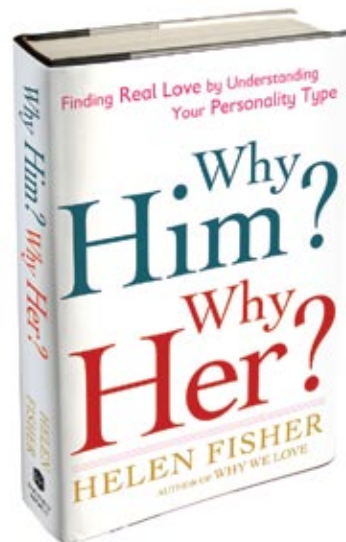
Wootton notes that the changes his team saw were linked with growing levels of atmospheric CO₂, but he readily acknowledges that the global-warming gas might not be the main culprit in this surge in acidity. Instead the acidification the researchers observed could have resulted from a nearby upwelling of deep ocean water loaded with carbon, so the results might not apply to the oceans as a whole. Still, the acidity readings along the Pacific coast of the U.S. and in the Netherlands do seem to be rising, Wootton says, “and that seems consistent with our pattern.” Marine life, it seems, may not have the luxury of time to act as a buffer against the changing waters.



CORALS and other marine life rich in calcium carbonate lose as the seas become acidic.

WHY LEAVE IT UP TO CUPID?

Discover the science behind human attraction—and the answer to love’s greatest question . . .



Helen Fisher, anthropologist and scientific adviser for Chemistry.com, applies the latest breakthroughs in neuroscience, brain chemistry, and evolutionary psychology to identify four personality types determined by your brain’s chemistry.

FIND OUT

- Why you fall in love with one person rather than another
- Which chemicals predispose your brain—and your heart
- Who is the best match for your personality type

Take Helen Fisher’s Personality Test at whyhimwhyher.com

Discover Who You Really Are and Who You Really Love

Also available on Macmillan Audio



Henry Holt and Company

PHYSICS

New Kind of Thermometer

For physicists, measuring temperature takes more than reading a column of mercury. They would like to define it in terms of a physical constant, just as length is described with respect to the speed of light (one meter is the distance traveled by light in an absolute vacuum in $1/299,792,458$ of one second). Currently the basic unit of temperature, the kelvin, is awkwardly defined as $1/273.16$ the difference between absolute zero and the triple point of water—that is, when water's gas, liquid and solid phases may co-exist at a certain pressure. One kelvin spans the same increment as one degree Celsius.

Now physicists have invented an electronic thermometer that ties temperature directly to a fundamental number—namely, the Boltzmann constant, a value related to the kinetic energy of molecules. (The constant is typically abbreviated in high school chemistry as k or k_B .) The device centers on the fact that in an array of tunnel junctions—thin, insulating layers sandwiched between electrodes—the electrical conductance can change in a manner directly proportional to the Boltzmann constant multiplied by the temperature.

Although such Coulomb blockade thermometry, as the technique is called, already appears in some specialized devices, fluctuations in the electronic properties of existing versions make them unreliable at very low temperatures. The new thermometer, made by scientists at the Helsinki University of Technology, works down to 150 millikelvins. Moreover, the Finnish physicists say it can be mass-produced using standard semiconductor manufacturing methods. They describe their work in the November 14 *Physical Review Letters*. —Charles Q. Choi



DIAGNOSTICS

Amnio Alternative

Amniocentesis and other prenatal tests designed to assess fetal health carry a small risk of miscarriage. Now Chinese researchers may have found an alternative diagnostic method based on a technique that distinguishes maternal DNA from fetal DNA in the mother's blood. That ability could lead to simple, no-risk blood tests that determine whether a fetus has a problem caused by single-gene mutations, such as cystic fibrosis and sickle cell anemia. The fetal DNA, which tends to be shorter than that of the mother, is duplicated and subjected to a “molecular counting” technique that tallies both mutant and normal genetic material. Researchers can use the data to determine whether the fetus has inherited a monogenetic disease. The San Diego-based biotech company Sequenom plans to develop the test for commercial distribution. The study appears in the December 16, 2008, *Proceedings of the National Academy of Sciences USA*. —Gary Stix

PRENATAL EXAMS to see if a fetus has inherited a genetic disorder could be replaced with no-risk blood tests.

Data Points
Truncated Lives

Zoo elephants live much shorter lives than their wild counterparts, according to a study based on some 4,500 elephants in European zoos and wildlife refuges. Infant mortality runs higher in captives, too—especially among Asian pachyderms, which suggests that something during gestation or early infancy raises the risk for the zoo-born. The data, however, may not reflect the latest zoo practices, which are more animal-friendly.

Median life span in years:

	In zoos	In the wild
African elephants	16.9	56
Asian elephants	18.9	41.7

Percent of infants born in first-time pregnancies that do not survive their first year:

	In zoos	In the wild
African elephants	25.9	18.7
Asian elephants	58.3	17.4



SOURCE: Science, December 12, 2008

NANOTECH

Booby Traps for Bacteria

Hollow capsules made of an organic conducting polymer could act as “roach motels” for bacteria. The microbes, which have an overall negative electrical charge, can get stuck on thin sheets or filaments extruding from the positively charged traps. When exposed to light, the capsules produce a very reactive form of oxygen highly toxic to bacteria—after one hour they killed more than 95 percent of nearby germs. The particles, built by scientists at the University of Florida and the University of New Mexico, can be applied to various surfaces, including medical equipment. The findings were presented online November 24 by *ACS Applied Materials & Interfaces*. —Charles Q. Choi

PRESCRIPTIONS

Rx Generics via Electronics

Doctors who prescribe electronically are more likely to select generics than pricey brand-name meds. In an 18-month study of more than 35,000 Massachusetts physicians, researchers found that an electronic system, which enables practitioners to tap in prescriptions and send them to pharmacies wirelessly, boosted the popularity of generics from 55 percent of all prescriptions to 61 percent. In contrast, a control group of physicians—who were not taught how to e-prescribe—were less inclined to go generic: those prescriptions increased from about 53 to 56 percent during the study period. So far only about 20 percent of physicians e-prescribe; if the practice were widely adopted, nearly \$4 million per 100,000 patients could be saved annually, according to the study authors. The findings are in the December 8, 2008, *Archives of Internal Medicine*. —Jordan Lite



LOWER TECH, HIGHER PRICES: Handwritten prescriptions are less likely to go generic.

BIOLOGY

Bug vs. Bug

Why can mosquitoes carry deadly viruses, such as West Nile and dengue, without succumbing to them? The prevailing theory maintained that the viruses and mosquitoes evolved to live in harmony. But entomologists have found quite the opposite to be true. They infected mosquitoes with a test virus and found that the

VIRAL VECTOR: Mosquito immune system keeps deadly viruses in check.

mosquito's immune system cut up the pathogen's genetic material so that the insect did not get sick. In contrast, when the mosquitoes were given a genetically modified version that blocked the gene-chopping mechanism, the insects could not mount an attack against the invader and died off more than four times as quickly. The discovery might lead to antivirals fashioned to mimic the mosquito's virus-killing tricks. The *Proceedings of the National Academy of Sciences USA* published the findings online December 1, 2008. —Susannah F. Locke



In Brief

SUBSURFACE GLACIERS ON MARS

Vast glaciers lie buried below thin layers of crustal debris on Mars, according to ground-penetrating radar from the Mars Reconnaissance Orbiter. Because current conditions on the Red Planet at the regions measured—between 30 and 60 degrees south latitude—do not support the development of ice, the glaciers probably took shape in the distant past, when Martian climate patterns were different. The debris covering protected the ice and kept it from sublimating into water vapor. The ice formations could constitute the largest stores of water on Mars outside its polar areas. —John Matson

EXOPLANETARY CARBON DIOXIDE

The Hubble Space Telescope has discovered carbon dioxide (CO₂) in the atmosphere of a planet outside our solar system. The exoplanet, called HD 189733b, is roughly the mass of Jupiter and orbits a star 63 light-years away. Scientists determined its atmospheric components by comparing the light spectrum from the star with that from the star and planet combined. Besides CO₂, the data reveal the existence of carbon monoxide, and previous findings indicate the presence of water vapor and methane. Although HD 189733b, which orbits very close to its parent star, is much too steamy for life as we know it, the finding shows that techniques exist to find markers of life on other worlds. —John Matson

WEAK ON THE NANO RISK

The plan of the National Nanotechnology Initiative (NNI) to ensure the safety of nanomaterials contains serious weaknesses, according to a December 10, 2008, report by the National Research Council. The NNI has a strategy to assess the risks of these substances, which include carbon nanotubes for strong materials and silver particles for antibacterial activity. But the council has found several flaws—for instance, the NNI has neither a summary of current safety knowledge nor an adequate way to hear from industry, academia and consumer advocates. The NNI says that it has begun pursuing some fixes but that it will need Congress to implement others. —Philip Yam



Read More ...

News Scan stories with this icon have extended coverage on www.SciAm.com/feb2009

SciAm Perspectives

A Scientific Stimulus for the U.S.

The right investments could help restore the nation's economic strength and environmental sustainability

BY THE EDITORS

One of the first orders of business for newly sworn-in President Barack Obama will be to push through a gigantic stimulus package to revive the U.S. economy from its coma. Debate swirls around how to spend that money; we would like to offer support for certain uses that seem both economically and scientifically worthy.

Obama has repeatedly emphasized, both on the campaign trail and in his postelection “fireside chats” on YouTube, that his economic recovery plan will steer massive funding toward America’s decaying and outmoded infrastructure. Multiple studies have documented how desperately U.S. bridges, highways, railways, dams, waterworks and other public resources need repair, modernization or replacement. In 2005 the American Society of Civil Engineers graded the state of U.S. infrastructure as “poor” and estimated that \$1.6 trillion would be needed over five years to fix it.

So the nation’s infrastructure could surely absorb abundant stimulus spending. Moreover, the money would be well spent as an investment: according to economist Mark Zandi of Moody’s Economy.com, a dollar spent on infrastructure typically adds \$1.59 to the gross domestic product.

Economic sustainability, however, is intertwined with energy and environmental sustainability, particularly the fight against global warming. Investments in our country’s transportation system should not merely repair it: where possible, they should strategically realign it to encourage more use of mass transportation instead of private automobiles to help the U.S. curb its energy consumption and carbon dioxide emissions.

To those ends, Obama has already pledged that his stimulus plan will make public buildings, including schools, more energy-efficient. Kudos on that good start, but don’t stop there, Mr. President. After decades of neglect as compared with heavily subsidized fossil fuels, every aspect of alternative energy development—from basic research to adoption and deployment—needs

more investment, too. With the recent collapse of oil prices during the global recession, alternative energy technologies could use financial shelter while they develop.

The most widely beneficial stimulus investment, however, could be for the government to subsidize the improvement of the electrical infrastructure of the U.S. The grassroots adoption of solar, wind, geothermal and other clean power is stifled by the grid’s inefficiency at handling highly variable inputs and transmitting power over long distances. An investment in a 21st-century grid would encourage energy reform without particularly favoring any specific production technology.

Universal health care reform—essential for ensuring that the benefits of biomedical research reach everyone—was a high priority early in this past presidential campaign season that deserves to be addressed again now. Even putting aside humanitarian and public health concerns, the economic case for overhauling our fragmented, employer-based health care delivery system is compelling. The catastrophe for Detroit’s automakers, for example, was certainly accelerated by heavy employee health care costs that more sensibly should be a public responsibility. Projections by the Congressional Budget Office indicate that in the decades to come, Medicaid and Medicare will be the biggest drivers of federal spending—far larger than defense, Social Security or other entitlements; health care costs overall will double their share of the national economy by 2050. Only prompt, comprehensive, fundamental reform can stop health care from thwarting any attempt to retire the nation’s bailout-and-stimulus debt.

Obama and Congress may come under great pressure to try to offset these colossal outlays (if only symbolically) by slashing other parts of the budget. We can appreciate the discipline of austerity but still hope that no one will be tempted to take an ax to the general research budget. Scientific discovery remains the fuel for technological innovation, and ultimately innovation will be what helps the U.S. salvage the economy.



Sustainable Developments

Transforming the Auto Industry

Only a partnership between the public and private sectors can help the Big Three roll into the future

BY JEFFREY D. SACHS



The U.S. auto industry has been widely vilified in recent months, with public opinion running strongly against government financial support for it. The industry fell into a trap of high costs, including unaffordable benefits and a morass of regulatory and contractual obligations that

enabled foreign producers to take a growing share of the U.S. market. Worse, the Big Three (Chrysler, Ford and General Motors) continued to promote gas-guzzling SUVs while the risks to the climate and U.S. energy security mounted. To some extent, the industry is also paying the price for spiraling national health costs, which should be under better public control, and for the U.S.'s inadequate fuel-efficiency policies and low gasoline taxes (in comparison with Europe and East Asia), which facilitated consumer demand for large vehicles. Yet many of the industry's problems indeed result from its own strategic miscalculations.

Still, the scorn for the industry misses four crucial points. First, a collapse of the Big Three would add another economic calamity to the crisis-roiled economy. Millions of jobs would be lost in places with very high unemployment and no offsetting job creation. Second, the automakers' plight is the result of the dramatic collapse of all domestic vehicle sales rather than the U.S.'s declining share of those sales. The Big Three would not be at the precipice of bankruptcy were it not for the worst recession since the Great Depression. Third, the public and political leadership bear huge co-responsibility with industry for the misguided SUV era, with its flagrant neglect of energy security, climate risks and unsustainable household borrowing.

Fourth, and most crucially, the changeover to high-mileage automobiles must be a public-private effort. Major technological change, such as from internal-combustion engines to electric vehicles recharged on a clean power grid or with hydrogen fuel cells, requires a massive infusion of public policy and public funding. Research and development depends on huge outlays, and many of the fruits of R&D should—and in any event will—become public goods rather than private intellectual property. That is why for a century the U.S. government has practiced public financing for R&D in many industries, including aviation,

computers, telephony, the Internet, drug development, advanced plant breeding, satellites and GPS.

To bemoan that the forthcoming Chevy Volt plug-in hybrid will have a first-year price tag of \$40,000 is to miss the point. Early-stage costs are inevitably far above those seen in the long run. Public policy should help promote this transition through such measures as the public-sector procurement of early models for official vehicle fleets, special tax and financing incentives for early purchasers, and higher gasoline taxes that internalize the costs of climate change and oil import dependence.

U.S. financing of sustainable energy technologies, such as for the high-performance batteries that are the limiting factor in high-performance plug-in hybrids, has been dreadfully small ever since President Ronald Reagan reversed the energy investments started by President Jimmy Carter. According to International

Energy Agency data, U.S. federal spending on all energy R&D amounted to just \$3 billion or so annually in recent years—less than two days of Pentagon spending. This neglect is finally changing with the pledge of \$25 billion in loans for technological upgrades approved in 2007 energy

legislation, but that money has not been mobilized and may arrive too late. Direct grants for R&D to government laboratories, companies and academia should be increased.

The U.S. needs a public-private technology policy, not merely finger-pointing at the private sector. GM's Chevy Volt, Chrysler's new Extended-Range Electric Vehicles and the large-scale efforts to produce a fuel-cell vehicle within a decade all require public backing, with R&D for basic technologies, support for early-stage demonstrations and dissemination, higher gas taxes to reflect security and climate costs, and public investments in complementary technologies, such as a clean power grid to charge the cars. It would be a mistake of historic proportions to let the industry die on the threshold of vital transformative change. ■

Jeffrey D. Sachs is director of the Earth Institute at Columbia University (www.earth.columbia.edu).



Skeptic

Darwin Misunderstood

On the 200th anniversary of Charles Darwin's birthday two myths persist about evolution and natural selection

BY MICHAEL SHERMER



On July 2, 1866, Alfred Russel Wallace, the co-discoverer of natural selection, wrote to Charles Darwin to lament how he had been “so repeatedly struck by the utter inability of numbers of intelligent persons to see clearly or at all, the self acting & necessary effects of *Nat Selection*,

that I am led to conclude that the term itself & your mode of illustrating it, however clear & beautiful to many of us are yet not the best adapted to impress it on the general *naturalist public*.” The source of the misunderstanding, Wallace continued, was the name itself, in that it implies “the constant watching of an intelligent ‘chooser’ like man’s selection to which you so often compare it,” and that “thought and direction are essential to the action of ‘Natural Selection.’” Wallace suggested redacting the term and adopting Herbert Spencer’s phrase “survival of the fittest.”

Unfortunately, that is what happened, and it led to two myths about evolution that persist today: that there is a prescient directionality to evolution and that survival depends entirely on cut-throat competitive fitness.

Contrary to the first myth, natural selection is a description of a process, not a force. No one is “selecting” organisms for survival in the benign sense of pigeon breeders selecting for desirable traits in show breeds or for extinction in the malignant sense of Nazis selecting prisoners at death camps. Natural selection is nonprescient—it cannot look forward to anticipate what changes are going to be needed for survival. When my daughter was young, I tried explaining evolution to her by using polar bears as an example of a “transitional species” between land mammals and marine mammals, but that was wrong. Polar bears are not “on their way” to becoming marine mammals. They are well adapted for their arctic environment.

Natural selection simply means that those individuals with variations better suited to their environment leave behind more offspring than individuals that are less well adapted. This outcome is known as “differential reproductive success.” It may be, as the second myth holds, that organisms that are bigger, stronger, faster and brutishly competitive will reproduce more successfully, but it is just as likely that organisms that are smaller,

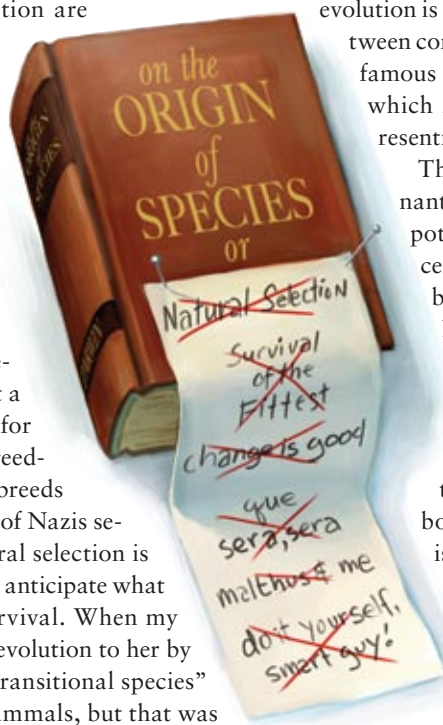
weaker, slower and socially cooperative will do so as well.

This second notion in particular makes evolution unpalatable for many people, because it covers the theory with a darkened patina reminiscent of Alfred, Lord Tennyson’s “nature, red in tooth and claw.” Thomas Henry Huxley, Darwin’s “bulldog” defender, promoted this “gladiatorial” view of life in a series of popular essays on nature “whereby the strongest, the swiftest, and the cunningest live to fight another day.” The myth persists. In his recent documentary film *Expelled: No Intelligence Allowed*, Ben Stein linked Darwinism to Communism, Fascism and the Holocaust. Former Enron CEO Jeff Skilling misread biologist Richard Dawkins’s book *The Selfish Gene* to mean that evolution is driven solely by ruthless competition, both between corporations and within Enron, leading to his infamous “rank and yank” employee evaluation system, which resulted in massive layoffs and competitive resentment.

This view of life need not have become the dominant one. In 1902 the Russian anarchist Petr Kropotkin published a rebuttal to Huxley and Spencer in his book *Mutual Aid*. Calling out Spencer by phrase, Kropotkin observed: “If we ... ask Nature: ‘who are the fittest: those who are continually at war with each other, or those who support one another?’ we at once see that those animals which acquire habits of mutual aid are undoubtedly the fittest.” Since that time science has revealed that species practice both mutual struggle and mutual aid. Darwinism, properly understood, gives us a dual disposition of selfishness and selflessness, competitiveness and cooperativeness.

Darwin was born on February 12, 1809, the same day as Abraham Lincoln, who also struggled to reconcile our binary natures in his first inaugural address on the eve of the Civil War: “The mystic chords of memory, stretching from every battlefield and patriot grave to every living heart and hearthstone all over this broad land, will yet swell the chorus of the Union, when again touched, as surely they will be, by the better angels of our nature.” ■

Michael Shermer is publisher of *Skeptic* (www.skeptic.com) and author of *Why Darwin Matters*.



Anti Gravity

Not a Close Shave

Is the main purpose of airport security to keep us unkempt?

BY STEVE MIRSKY



"'Curiouser and curiouser!' cried Alice (she was so much surprised, that for the moment she quite forgot how to speak good English). 'Now I'm opening out like the largest telescope that ever was! Good-by, feet!... Oh, my poor little feet, I wonder who will put on your shoes and stockings for you now, dears?'" The smart money says that it won't be the folks from the Transportation Security Administration, who make two million travelers take their shoes off every day at airports in the U.S.

Lewis Carroll's Alice would have had trouble distinguishing reality from Wonderland had she been with me on the Sunday after Thanksgiving as I watched a TSA officer confiscate my father's aftershave at the airport in Burlington, Vt. It was a 3.25-ounce bottle, clearly in violation of the currently permissible three-ounce limit for liquids. Also clear was the bottle, which was obviously only about a quarter full. So even the members of some isolated human populations that have never developed sophisticated systems for counting could have determined that the total amount of liquid in the vessel was far less than the arbitrarily standardized three ounces. But the TSA guy took the aftershave, citing his responsibility to go by the volume listed on the label. (By the way, the three-ounce rule is expected to be phased out late in 2009. Why not tomorrow? Because of the 300-day-rules-change rule, which I just made up.)

Feeling curiouser, I did a gedankenexperiment: What if the bottle had been completely empty—would he have taken it then? No, I decided. When empty, the bottle becomes just some plastic in a rather mundane topological configuration. Not to mention that if you really banned everything with the potential to hold more than three ounces of liquid, you couldn't let me have my shoes back. You also couldn't allow me to bring my hands onboard. I kept these thoughts to myself, of course, because I wanted to fly home, not spend the rest of the day locked in a security office explaining what a gedankenexperiment was.

I first commented on what I used to call "the illusion of security" in this space in July 2003, after attending a conference on freedom and privacy. We heard the

story of an airline pilot who had his nail clippers snatched away by the TSA just before boarding his plane. He then walked into a cockpit equipped with an ax. (Which is a horrible tool for cutting your nails, although, I have to admit, my dad might try. A former U.S. Marine and builder, he does his manicuring with a foot-long metal carpenter's file and some 80-grit sandpaper. And you wonder how I got to be this way.)

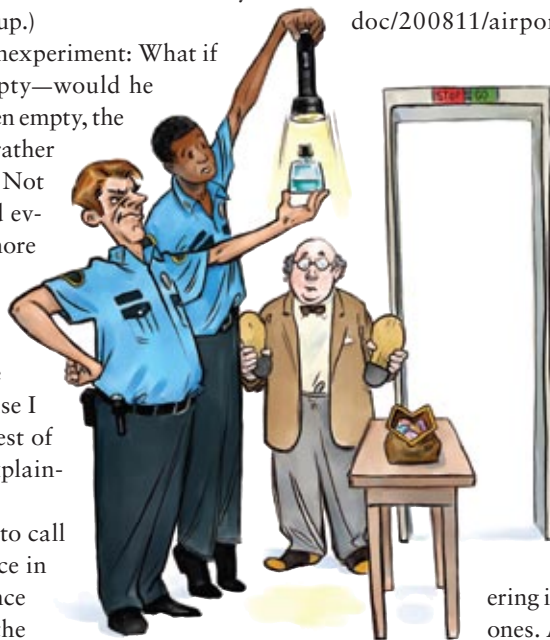
It used to be that you could bring shaving cream with you when boarding a plane, but they would take away your razor. Now you can carry on a razor, and they take away your shaving cream. (They did indeed seize my dad's shaving cream at the airport in Fort Lauderdale the Monday before Thanksgiving.)

Although the mostest curiouser thing has to be when hundreds of people docilely snake through security lines amid announcements that the "threat of a terrorist attack is high." Compared to what? The day before, perhaps, when the real threat posed by terrorists to your life was much, much smaller than your chances of dying in the bathtub. And today the threat is only much smaller than your chances of dying in the bathtub. Here's how you know that the terrorist threat isn't really high: the airport is still open, and your flight hasn't been canceled.

A much better term than "illusion of security" can be found in an article by Jeffrey Goldberg in the November 2008 issue of the *Atlantic*: "security theater" (www.theatlantic.com/doc/200811/airport-security). Goldberg holds that TSA

agents and passengers go through performances designed to make everybody feel better, but with little effect. He talks about how he has been able to carry knives and box cutters onto planes—he even got past security with a device on his torso called a Beerbelly, a bladder that holds up to 80 ounces of liquid you can drink from through a tube.

Goldberg didn't fill the thing up, but he did exceed the three-ounce limit by just 21 ounces. He believes that our current airport procedures may succeed in catching dumb terrorists. But the time, energy and money would be better spent on gathering intelligence if we want to catch the smart ones. And keep my dad clean-shaven. ■



Naked Singularities

The black hole has a troublesome sibling, the naked singularity. Physicists have long thought—hoped—it could never exist. But could it?

BY PANKAJ S. JOSHI

KEY CONCEPTS

- Conventional wisdom has it that a large star eventually collapses to a black hole, but some theoretical models suggest it might instead become a so-called naked singularity [see box on opposite page for definition]. Sorting out what happens is one of the most important unresolved problems in astrophysics.
- The discovery of naked singularities would transform the search for a unified theory of physics, not least by providing direct observational tests of such a theory.

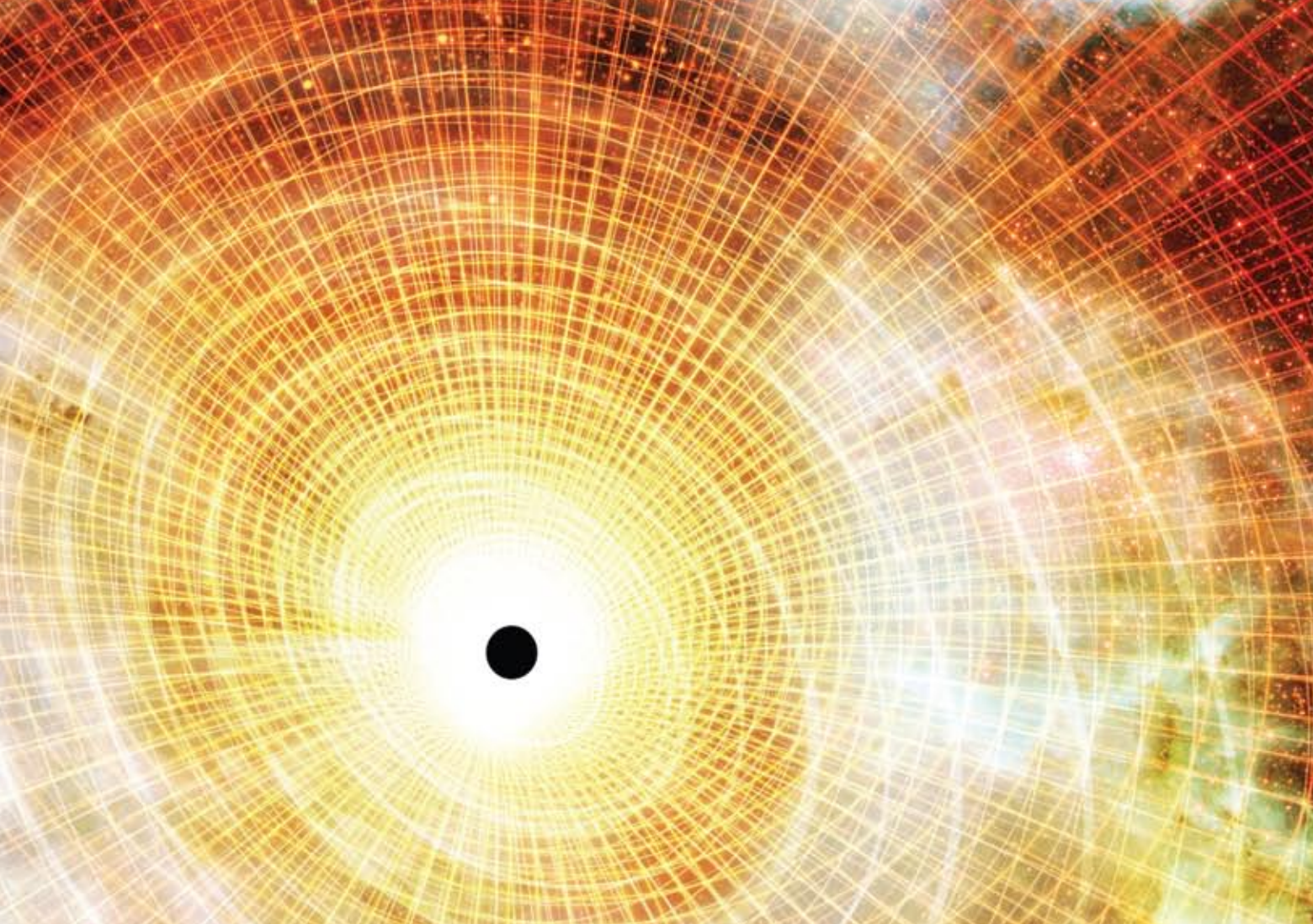
—The Editors

Modern science has introduced the world to plenty of strange ideas, but surely one of the strangest is the fate of a massive star that has reached the end of its life. Having exhausted the fuel that sustained it for millions of years, the star is no longer able to hold itself up under its own weight, and it starts collapsing catastrophically. Modest stars like the sun also collapse, but they stabilize again at a smaller size. Whereas if a star is massive enough, its gravity overwhelms all the forces that might halt the collapse. From a size of millions of kilometers across, the star crumples to a pinprick smaller than the dot on an “i.”

Most physicists and astronomers think the

result is a black hole, a body with such intense gravity that nothing can escape from its immediate vicinity. A black hole has two parts. At its core is a singularity, the infinitesimal point into which all the matter of the star gets crushed. Surrounding the singularity is the region of space from which escape is impossible, the perimeter of which is called the event horizon. Once something enters the event horizon, it loses all hope of exiting. Whatever light the falling body gives off is trapped, too, so an outside observer never sees it again. It ultimately crashes into the singularity.

But is this picture really true? The known laws of physics are clear that a singularity forms,



but they are hazy about the event horizon. Most physicists operate under the assumption that a horizon must indeed form, if only because the horizon is very appealing as a scientific fig leaf. Physicists have yet to figure out what exactly happens at a singularity: matter is crushed, but what becomes of it then? The event horizon, by hiding the singularity, isolates this gap in our knowledge. All kinds of processes unknown to science may occur at the singularity, yet they have no effect on the outside world. Astronomers plotting the orbits of planets and stars can safely ignore the uncertainties introduced by singularities and apply the standard laws of physics with confidence. Whatever happens in a black hole stays in a black hole.

Yet a growing body of research calls this working assumption into question. Researchers have found a wide variety of stellar collapse scenarios in which an event horizon does not in fact form, so that the singularity remains exposed to our view. Physicists call it a naked singularity. Matter and radiation can both fall in and come out. Whereas visiting the singularity inside a black hole would be a one-way trip, you could in

principle come as close as you like to a naked singularity and return to tell the tale.

If naked singularities exist, the implications would be enormous and would touch on nearly every aspect of astrophysics and fundamental physics. The lack of horizons could mean that mysterious processes occurring near the singularities would impinge on the outside world. Naked singularities might account for unexplained high-energy phenomena that astronomers have seen, and they might offer a laboratory to explore the fabric of spacetime on its finest scales.

Cosmic Censor

Event horizons were supposed to have been the easy part about black holes. Singularities are clearly mysterious. They are places where the strength of gravity becomes infinite and the known laws of physics break down. According to physicists' current understanding of gravity, encapsulated in Einstein's general theory of relativity, singularities inevitably arise during the collapse of a giant star. General relativity does not account for the quantum effects that become important for microscopic objects, and those

CLOTHED OR NAKED?

Black holes and naked singularities are two possible outcomes of the collapse of a dying massive star. At the heart of each is a singularity—a wad of matter so dense that it requires new laws of physics to describe. Anything that hits the singularity gets destroyed.

In a black hole, the singularity is “clothed”—that is, surrounded by a boundary called the event horizon that hides it. Nothing that falls through this surface can ever get back out.

A naked singularity has no such boundary. It is visible to outside observers, and objects that fall toward the singularity can in principle reverse course right up to the moment of impact.

effects presumably intervene to prevent the strength of gravity from becoming truly infinite. But physicists are still struggling to develop the quantum theory of gravity they need to explain singularities.

By comparison, what happens to the region of spacetime around the singularity seems as though it should be rather straightforward. Stellar event horizons are many kilometers in size, far larger than the typical scale of quantum effects. Assuming that no new forces of nature intervene, horizons should be governed purely by general relativity, a theory that is based on well-understood principles and has passed 90 years of observational tests.

That said, applying the theory to stellar collapse is still a formidable task. Einstein's equations of gravity are notoriously complex, and solving them requires physicists to make simplifying assumptions. American physicists J. Robert Oppenheimer and Hartland S. Snyder—and, independently, Indian physicist B. Datt—made an initial attempt in the late 1930s. To simplify the equations, they considered only perfectly spherical stars, assumed the stars consisted of gas of a homogeneous (uniform) density and neglected gas pressure. They found that as this idealized star collapses, the gravity at its surface intensifies and eventually becomes strong enough to trap all light and matter, thereby forming an event horizon. The star becomes invisible to outside observers and soon thereafter collapses all the way down to a singularity.

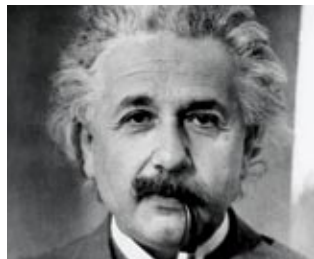
Real stars, of course, are more complicated. Their density is inhomogeneous, the gas in them exerts pressure, and they can assume other shapes. Does every sufficiently massive collapsing star turn into a black hole? In 1969 University of Oxford physicist Roger Penrose suggested that the answer is yes. He conjectured that the formation of a singularity during stellar collapse necessarily entails the formation of an event horizon. Nature thus forbids us from ever seeing a singularity, because a horizon always cloaks it. Penrose's conjecture is termed the cosmic censorship hypothesis. It is only a conjecture, but it underpins the modern study of black holes. Physicists hoped we would be able to prove it with the same mathematical rigor we used to show the inevitability of singularities.

Singularities au Naturel

That has not happened. Instead of coming up with a direct proof of censorship that applies under all conditions, we have had to embark on

FATHER FIGURES

The ongoing debate over whether naked singularities can form is one part of a broader story arc about black holes.



The general theory of relativity predicted black holes, but Einstein doubted they could ever really form.



J. Robert Oppenheimer (later head of the Manhattan Project) and others showed they could.



Stephen Hawking and Roger Penrose (below) proved that singularities were, in fact, unavoidable.



Penrose conjectured that singularities had to remain clothed by event horizons. Others disagree.

the longer route of analyzing case studies of gravitational collapse one by one, gradually embellishing our theoretical models with the features that the initial efforts lacked. In 1973 German physicist Hans Jürgen Seifert and his colleagues considered inhomogeneity. Intriguingly, they found that layers of infalling matter could intersect to create momentary singularities that were not covered by horizons. But singularities come in various types, and these ones were fairly benign. Although the density at one location became infinite, the strength of gravity did not, so the singularity did not crush matter and infalling objects to an infinitesimal pinprick. Thus, general relativity never broke down, and matter continued to move through this location rather than meeting its end.

In 1979 Douglas M. Eardley of the University of California, Santa Barbara, and Larry Smarr of the University of Illinois at Urbana-Champaign went a step further and performed a numerical simulation of a star with a realistic density profile: highest at its center and slowly decreasing toward the surface. An exact paper-and-pencil treatment of the same situation, undertaken by Demetrios Christodoulou of the Swiss Federal Institute of Technology in Zurich, followed in 1984. Both studies found that the star shrank to zero size and that a naked singularity resulted. But the model still neglected pressure, and Richard P.A.C. Newman, then at the University of York in England, showed that the singularity was again gravitationally weak.

Inspired by these findings, many researchers, including me, tried to formulate a rigorous theorem that naked singularities would always be weak. We were unsuccessful. The reason soon became clear: naked singularities are not always weak. We found scenarios of inhomogeneous collapse that led to singularities where gravity was strong—that is, genuine singularities that could crush matter into oblivion—yet remained visible to external observers. A general analysis of stellar collapse in the absence of gas pressure, developed in 1993 by Indresh Dwivedi, then at Agra University, and me, clarified and settled these points.

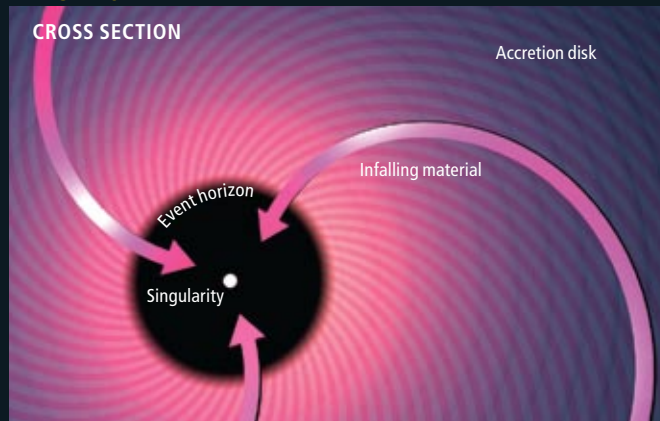
In the early 1990s physicists considered the effects of gas pressure. Amos Ori of the Technion-Israel Institute of Technology and Tsvi Piran of the Hebrew University of Jerusalem conducted numerical simulations, and my group solved the relevant equations exactly. Stars with a fully realistic relation between density and pressure could collapse to naked singularities.

[BASIC PRINCIPLES]

Two Monsters of the Cosmos

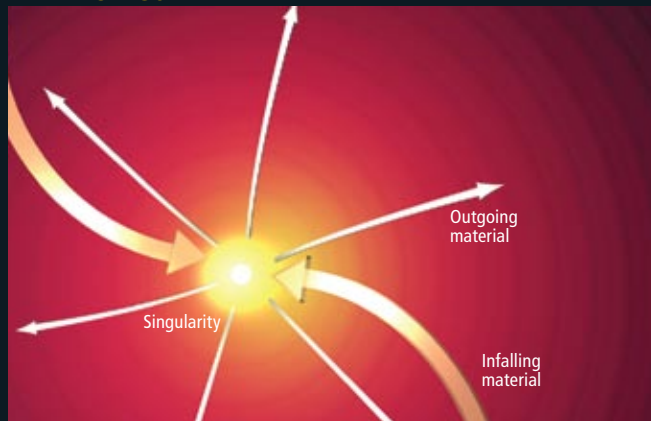
A naked singularity is essentially a black hole without the “black” part. It can both suck in and spit out matter and radiation. Accordingly, it would look different and have different effects on its immediate surroundings.

BLACK HOLE

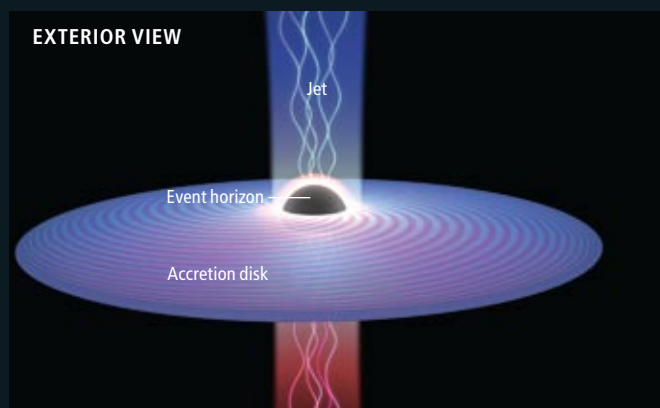


The key feature of a black hole is its event horizon, a surface that material can enter but not exit. Often the horizon is surrounded by a swirling, gaseous disk.

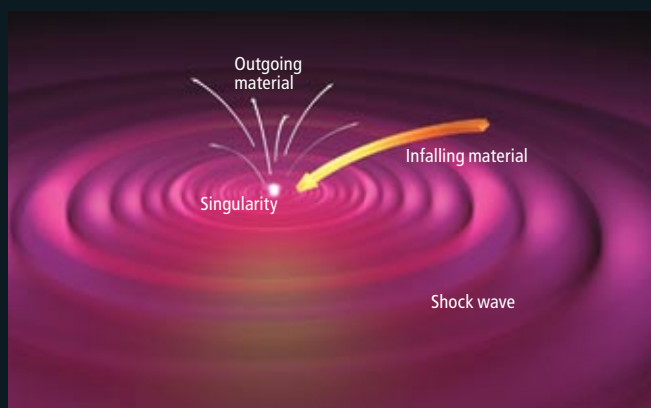
NAKED SINGULARITY



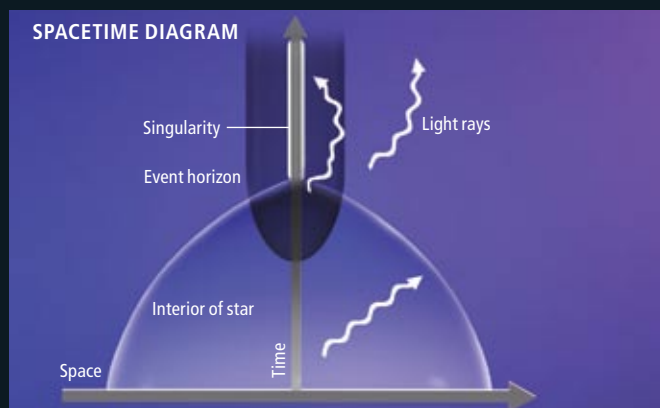
A naked singularity lacks an event horizon. Like a black hole, it can suck in material; unlike a black hole, it can also spit it out.



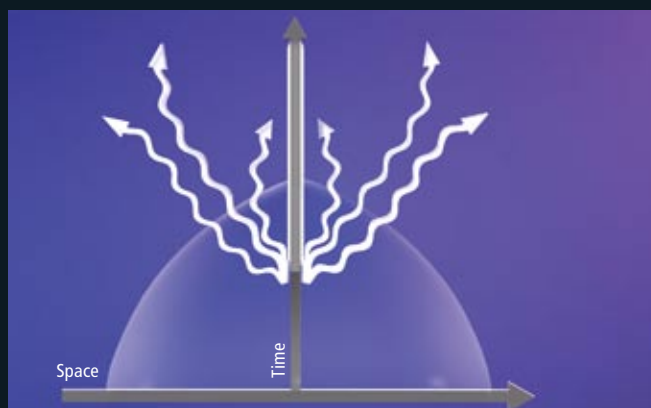
On the outside, the black hole looks like a pitch-black ball. The singularity lies within and cannot be seen. Friction within the surrounding disk generates intense radiation. Some disk material flies out in a jet; other material falls in.



A naked singularity looks like a tiny dust grain, except that it is almost inconceivably dense. Infalling matter can be seen all the way down to its final impact with the singularity. The intense gravity can generate powerful shock waves.



A homogeneous star without gas pressure collapses to a black hole. The star's gravity intensifies, bending the paths of moving objects (including light rays) by an increasing amount and eventually trapping them.



When the star is inhomogeneous, its gravity might never intensify enough to bend light rays back on themselves. The star collapses to a singularity, but the singularity is visible. [For step-by-step simulations, see box on next two pages.]

Two Ways to Crush a Star

Computer simulations reveal the circumstances under which a star collapses to a black hole or to a naked singularity. The simulations shown here treat the

star as a swarm of grains whose gravity is so powerful that other forces of nature, such as gas pressure, are unimportant.

BLACK HOLE



At about the same time, teams led by Giulio Magli of the Polytechnic University of Milan and by Kenichi Nakao of Osaka City University considered a form of pressure generated by rotation of particles within a collapsing star. They, too, showed that in a wide variety of situations, collapse ends in a naked singularity after all.

These studies analyzed perfectly spherical stars, which is not as severe a limitation as it might appear, because most stars in nature are very close to this shape. Moreover, spherical stars have, if anything, more favorable conditions for horizon formation than stars of other shapes do, so if cosmic censorship fails even for them, its prospects look questionable. That said, physicists have been exploring nonspherical collapse. In 1991 Stuart L. Shapiro of the University of Illinois and Saul A. Teukolsky of Cornell University presented numerical simulations in which oblong stars could collapse to a naked singularity. A few years later Andrzej Królak of the Polish Academy of Sciences and I studied nonspherical collapse and also found naked singularities. To be sure, both these studies neglected gas pressure.

Skeptics have wondered whether these situations are contrived. Would a slight change to the initial configuration of the star abruptly cause

an event horizon to cover the singularity? If so, then the naked singularity might be an artifact of the approximations used in the calculations and would not truly arise in nature. Some scenarios involving unusual forms of matter are indeed very sensitive. But our results so far also show that most naked singularities are stable to small variations of the initial setup. Thus, these situations appear to be what physicists call generic—that is, they are not contrived.

How to Beat the Censor

These counterexamples to Penrose's conjecture suggest that cosmic censorship is not a general rule. Physicists cannot say, "Collapse of any massive star makes a black hole only," or "Any physically realistic collapse ends in a black hole." Some scenarios lead to a black hole and others to a naked singularity. In some models, the singularity is visible only temporarily, and an event horizon eventually forms to cloak it. In others, the singularity remains visible forever. Typically the naked singularity develops in the geometric center of collapse, but it does not always do so, and even when it does, it can also spread to other regions. Nakedness also comes in degrees: an event horizon might hide the singularity from the prying eyes of faraway observers, whereas

NAKED SINGULARITY



1 This star has the shape of an American football.



2 It collapses toward its axis.



3 The star begins to look like a narrow spindle.



4 Although the gravity intensifies, it never becomes strong enough to trap light and form a horizon.

5 The density is highest near the two ends of the spindle, and singularities will form there.



6 No horizon ever emerges to conceal the singularities, so they remain visible to outside observers.



observers who fell through the event horizon could see the singularity prior to hitting it. The variety of outcomes is bewildering.

My colleagues and I have isolated various features of these scenarios that cause an event horizon to arise or not. In particular, we have examined the role of inhomogeneities and gas pressure. According to Einstein's theory, gravity is a complex phenomenon involving not only a force of attraction but also effects such as shearing, in which different layers of material are shifted laterally in opposite directions. If the density of a collapsing star is very high—so high that by all rights it should trap light—but also inhomogeneous, those other effects may create escape routes. Shearing of material close to a singularity, for example, can set off powerful shock waves that eject matter and light—in essence, a gravitational typhoon that disrupts the formation of an event horizon.

To be specific, consider a homogeneous star, neglecting gas pressure. (Pressure alters the details but not the broad outlines of what happens.) As the star collapses, gravity increases in strength and bends the paths of moving objects ever more severely. Light rays, too, become bent, and there comes a time when the bending is so severe that light can no longer propagate away

OBSERVING NAKED SINGULARITIES

Naked singularities could reveal their presence to astronomers in several ways:

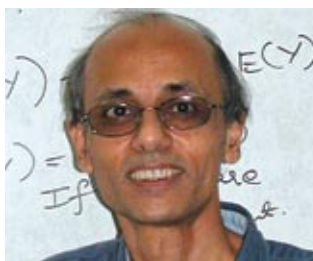
- High-energy explosions that produce naked singularities would brighten and fade in a distinctive way.
- Certain classes of gamma-ray bursts remain unexplained; naked singularities might account for them.
- Naked singularities would bend the light of background stars differently than black holes do.
- If a suspected black hole is spinning faster than a certain rate (which depends on its mass), it must actually be a naked singularity. The planned Square Kilometer Array (SKA) radio telescope would have the requisite precision to tell.

from the star. The region where light becomes trapped starts off small, grows and eventually reaches a stable size proportional to the star's mass. Meanwhile because the star's density is uniform in space and varies only in time, the entire star is crushed to a point simultaneously. The trapping of light occurs well before this moment, so the singularity remains hidden.

Now consider the same situation except that the density decreases with distance from the center. In effect, the star has an onionlike structure of concentric shells of matter. The strength of gravity acting on each shell depends on the average density of matter interior to that shell. Because the denser inner shells feel a stronger pull of gravity, they collapse faster than the outer ones. The entire star does not collapse to a singularity simultaneously. The innermost shells collapse first, and then the outer shells pile on, one by one.

The resulting delay can postpone the formation of an event horizon. If the horizon can form anywhere, it will form in the dense inner shells. But if density decreases with distance too rapidly, these shells may not constitute enough mass to trap light. The singularity, when it forms, will be naked. Therefore, there is a threshold: if the degree of inhomogeneity is very small, below a crit-

[THE AUTHOR]



Pankaj S. Joshi is a physics professor at the Tata Institute of Fundamental Research in Mumbai. He specializes in gravitation and cosmology. When not doing physics, he enjoys taking long walks in nature and listening to Indian and Western classical music. Joshi lives very close to the sites of the November 2008 terrorist attacks and was at home working on this article when they happened.

ical limit, a black hole will form; with sufficient inhomogeneity, a naked singularity arises.

In other scenarios, the salient issue is the rapidity of collapse. This effect comes out very clearly in models where stellar gas has converted fully to radiation and, in effect, the star becomes a giant fireball—a scenario first considered by Indian physicist P. C. Vaidya in the 1940s in the context of modeling a radiating star. Again there is a threshold: slowly collapsing fireballs become black holes, but if a fireball collapses rapidly enough, light does not become trapped and the singularity is naked.

Slimy

One reason it has taken so long for physicists to accept the possibility of naked singularities is that they raise a number of conceptual puzzles. A commonly cited concern is that such singularities would make nature inherently unpredictable. Because general relativity breaks down at singularities, it cannot predict what those singularities will do. John Earman of the University of Pittsburgh memorably suggested that green slime and lost socks could emerge from them. They are places of magic, where science fails.

As long as singularities remain safely ensconced within event horizons, this randomness remains contained and general relativity is a fully predictive theory, at least outside the horizon. But if singularities can be naked, their unpredictability would infect the rest of the universe. For example, when physicists applied general relativ-

ity to Earth's orbit around the sun, they would in effect have to make allowance for the possibility that a singularity somewhere in the universe could emit a random gravitational pulse and send our planet flying off into deep space.

Yet this worry is misplaced. Unpredictability is actually common in general relativity—and not always directly related to censorship violation. The theory permits time travel, which could produce causal loops with unforeseeable outcomes, and even ordinary black holes can become unpredictable. For example, if we drop an electric charge into an uncharged black hole, the shape of spacetime around the hole radically changes and is no longer predictable. A similar situation holds when the black hole is rotating. Specifically, what happens is that spacetime no longer neatly separates into space and time, so physicists cannot consider how the black hole evolves from some initial time into the future. Only the purest of pure black holes, with no charge or rotation at all, is fully predictable.

The loss of predictability and other problems with black holes actually stem from the occurrence of singularities; it does not matter whether they are hidden or not. The solution to these problems probably lies in a quantum theory of gravity, which will go beyond general relativity and offer a full explication of singularities. Within that theory, every singularity would prove to have a high but finite density. A naked singularity would be a “quantum star,” a hyperdense body governed by the rules of quantum gravity. What seems random would have a logical explanation.

Another possibility is that singularities may really have an infinite density after all—that they are not things to be explained away by quantum gravity but to be accepted as they are. The breakdown of general relativity at such a location may not be a failure of the theory per se but a sign that space and time have an edge. The singularity marks the place where the physical world ends. We should think of it as an event rather than an object, a moment when collapsing matter reaches the edge and ceases to be, like the big bang in reverse.

In that case, questions such as what will come out of a naked singularity are not really meaningful; there is nothing to come out of, because the singularity is just a moment in time. What we see from a distance is not the singularity itself but the processes that occur in the extreme conditions of matter near this event, such as shock waves caused by inhomogeneities in

[A RELATED PROBLEM]

Can a Black Hole Be Cracked Open?

Besides the collapse of a star, another way to create a naked singularity might be to take an existing black hole and destroy it. Although that sounds like an impossible, not to mention dangerous, task, the equations of general relativity show that an event horizon can exist only if the hole does not rotate too fast or does not have too much electrical charge. Most physicists think that the hole resists any efforts to spin or charge it up beyond the prescribed limits. But some think the hole might eventually succumb, causing the horizon to dissipate and expose the singularity.

Spinning up a black hole is not terribly difficult. Matter naturally

falls in with angular momentum, which drives the hole to spin faster, like pushing on a revolving door. Charging up a black hole is harder because a charged hole repels particles of the same charge and draws in oppositely charged ones, thereby neutralizing itself. But a vigorous infall of matter could overcome this tendency. The fundamental characteristic of a black hole—namely, its trait of gobbling up the matter around it that keeps it growing—could become the cause of its own destruction. Researchers continue to debate whether the black hole would eventually save itself or would crack open to reveal the singularity within. —P.S.J.

this ultradense medium or quantum-gravitational effects in its vicinity.

In addition to unpredictability, a second issue troubles many physicists. Having provisionally assumed that the censorship conjecture holds, they have spent the past several decades formulating various laws that black holes should obey, and these laws have the ring of deep truths. But the laws are not free of major paradoxes. For example, they hold that a black hole swallows and destroys information—which appears to contradict the basic principles of quantum theory [see “Black Holes and the Information Paradox,” by Leonard Susskind; *SCIENTIFIC AMERICAN*, April 1997]. This paradox and other predicaments stem from the presence of an event horizon. If the horizon goes away, these problems might go away, too. For instance, if the star could radiate away most of its mass in the late stages of collapse, it would destroy no information and leave behind no singularity. In that case, it would not take a quantum theory of gravity to explain singularities; general relativity might do the trick itself.

A Lab for Quantum Gravity

Far from considering naked singularities a problem, physicists can see them as an asset. If the singularities that formed in the gravitational collapse of a massive star are visible to external observers, they could provide a laboratory to study quantum-gravitational effects. Quantum gravity theories in the making, such as string theory and loop quantum gravity, are badly in need of some kind of observational input, without which it is nearly impossible to constrain the plethora of possibilities. Physicists commonly seek that input in the early universe, when conditions were so extreme that quantum-gravitational effects dominated. But the big bang was a unique event. If singularities could be naked, they would allow astronomers to observe the equivalent of a big bang every time a massive star in the universe ends its life.

To explore how naked singularities might provide a glimpse into otherwise unobservable phenomena, we recently simulated how a star collapses to a naked singularity, taking into account the effects predicted by loop quantum gravity. According to this theory, space consists of tiny atoms, which become conspicuous when matter becomes sufficiently dense; the result is an extremely powerful repulsive force that prevents the density from ever becoming infinite [see “Follow the Bouncing Universe,” by Mar-

Varieties of Stellar Collapse

Like humans, stars have a life cycle. They are born in gigantic clouds of dust and galactic material in the depths of space, evolve and shine for millions of years, and eventually enter a phase of dissolution and extinction. Stars shine by burning their nuclear fuel, which is initially mainly hydrogen; they fuse it into helium and, later, heavier elements. Each star attains a balance between the force of gravity, which pulls matter toward the center, and the outward pressures generated by fusion. This balance keeps the star stable—until all the fuel is converted to iron, which, in nuclear terms, is inert. Fusion then ceases, the all-pervasive force of gravity asserts itself, and the star begins to contract.

When our sun runs out of fuel, its core will contract under its own gravity until it is no bigger than Earth, at which point it will be supported by the force exerted by fast-moving electrons, called electron degeneracy pressure. The resulting object is called a white dwarf. Stars three to five times the mass of the sun settle to a different final state, a neutron star, in which gravity is so strong that atoms buckle, too. Supported by the pressure not of electrons but of neutrons, a neutron star is barely 10 kilometers in size.

Still more massive stars cannot settle either to a white dwarf or to a neutron star, because these forms of pressure are just not sufficient. Unless some other, unknown form of pressure comes into play, gravitational collapse becomes unstoppable. Gravity now is the sole operative force, and the final fate of the star is determined by Einstein's theory of gravitation. The theory indicates that the outcome is a singularity, and the question is whether this singularity is visible or not.

—P.S.J.

MORE TO EXPLORE

Black Holes, Naked Singularities and Cosmic Censorship. Stuart L. Shapiro and Saul A. Teukolsky in *American Scientist*, Vol. 79, No. 4, pages 330–343; July/August 1991.

Black Holes and Time Warps. Kip S. Thorne. W. W. Norton, 1994.

Bangs, Crunches, Whimpers, and Shrieks: Singularities and Acausalities in Relativistic Spacetimes. John Earman. Oxford University Press, 1995.

Quantum Evaporation of a Naked Singularity. Rituparno Goswami, Pankaj S. Joshi and Parampreet Singh in *Physical Review Letters*, Vol. 96, No. 3, Paper No. 031302; January 27, 2006. Available online at <http://arxiv.org/abs/gr-qc/0506129>

Gravitational Collapse and Spacetime Singularities. Pankaj S. Joshi. Cambridge University Press, 2007.

tin Bojowald; *SCIENTIFIC AMERICAN*, October 2008]. In our model, such a repulsive force dispersed the star and dissolved the singularity. Nearly a quarter of the mass of the star was ejected within the final fraction of a microsecond. Just before it did so, a faraway observer would have seen a sudden dip in the intensity of radiation from the collapsing star—a direct result of quantum-gravitational effects.

The explosion would have unleashed high-energy gamma rays, cosmic rays and other particles such as neutrinos. Upcoming experiments such as the Extreme Universe Space Observatory, a module for the International Space Station expected to be operational in 2013, may have the needed sensitivity to see this emission. Because the details of the outpouring depend on the specifics of the quantum gravity theory, observations would provide a way to discriminate among theories.

Either proving or disproving cosmic censorship would create a mini explosion of its own within physics, because naked singularities touch on so many deep aspects of current theories. What comes out unambiguously from the theoretical work so far is that censorship does not hold in an unqualified form, as it is sometimes taken to be. Singularities are clothed only if the conditions are suitable. The question remains whether these conditions could ever arise in nature. If they can, then physicists will surely come to love what they once feared. ■

Nanomedicine

Targets



CANCER

Viewing each human body as a system of interacting molecular networks and targeting disruptions in the system with nanoscale technologies can transform how disease is understood, attacked and possibly prevented

By James R. Heath, Mark E. Davis and Leroy Hood

KEY CONCEPTS

- A “systems” approach to medicine views the body as a complex network of molecular interactions that can be measured and modeled, revealing causes of disease such as cancer.
- Extremely miniaturized tools can inexpensively measure and manipulate molecules for systems medicine.
- Nanoscale therapies deliver precisely targeted treatments to tumors while avoiding healthy tissues.

—The Editors

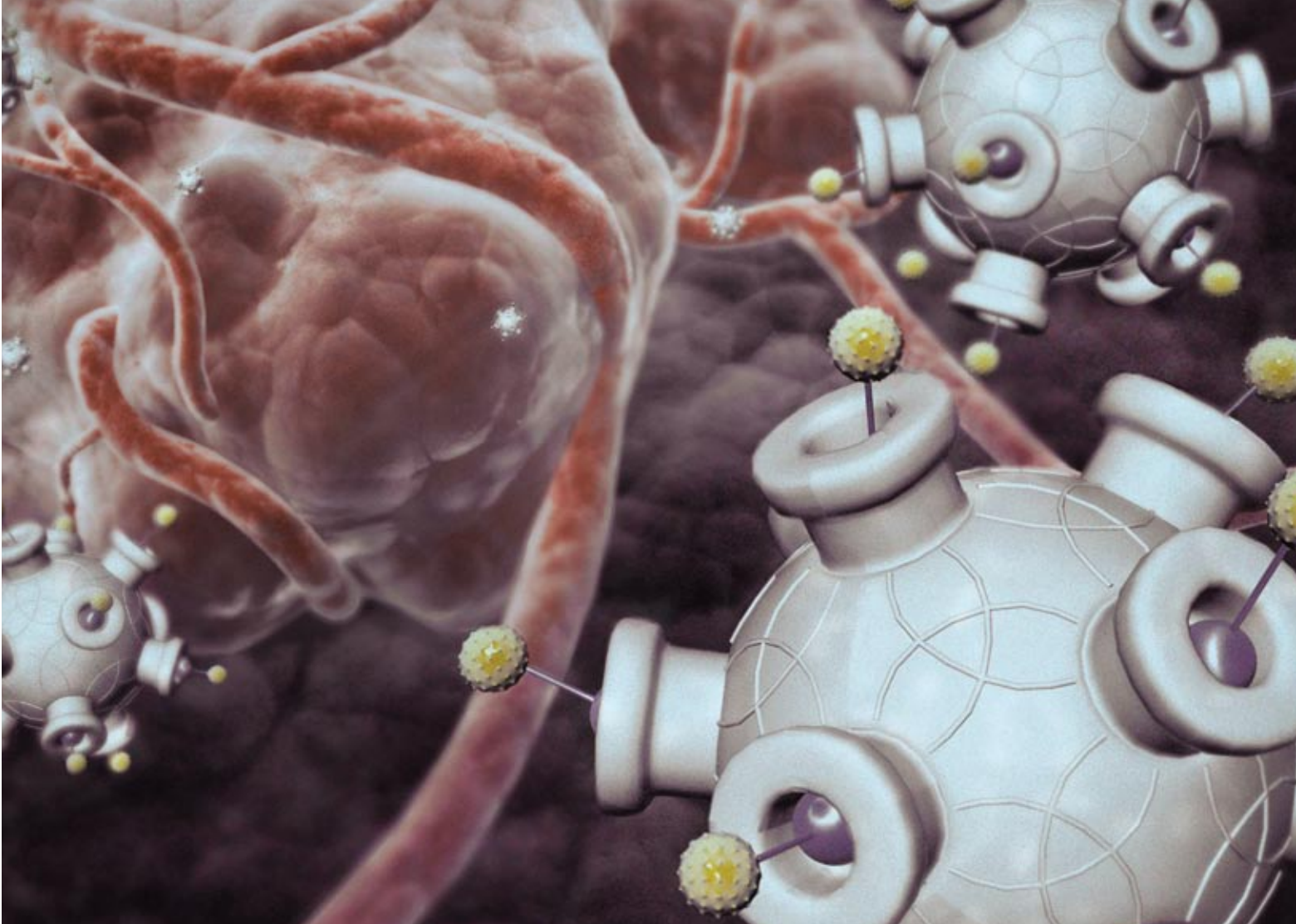
Before going to the gym for a workout or after indulging in cake at the office party, people with diabetes can use a portable monitor to take a quick blood glucose measurement and adjust their food or insulin intake to prevent extreme dips or spikes in blood sugar. The inexpensive finger-prick testing devices that allow diabetics to check their glucose levels throughout the day may sound like small conveniences. That is unless you are diabetic and can remember back a decade or more, when having that disease came with far more fear and guessing and far less control over your own well-being.

The quality of life afforded to diabetics by technologies that easily and inexpensively extract information from the body offers a glimpse of what all medicine could be like: more predictive and preventive, more personalized to the individual’s needs and enabling more participation in maintaining one’s own health. In fact, we believe that medicine is already headed in that direction, largely because of new technologies that make it possible to acquire and analyze biological information quickly and cheaply.

One of the keys to this evolution in medicine

is the extreme miniaturization of technologies for making diagnostic measurements from minuscule amounts of blood or even single cells taken from diseased tissues. These emerging tools, constructed at the scale of microns and nanometers (billionths of a meter), can manipulate and measure large numbers of biological molecules rapidly, precisely and, eventually, at a cost of pennies or less per measurement. That combination of cost and performance opens up new avenues for studying and treating disease by permitting the human body to be viewed as a dynamic system of molecular interactions. Such systems-level measurements are then integrated into computational models, which, in turn, can reveal early indicators of a problem. When these insights are combined with new nanotechnology-based therapies, the treatment can be targeted to the problem and only the problem, thereby avoiding serious side effects.

Although we anticipate that all medicine will eventually operate by these principles, cancer research offers current examples of how technology at the ultrasmall scale is providing the data needed to paint a big-picture systems view of disease.



Systems Medicine

Modeling a system requires vast amounts of data, and living organisms are full of information that could be described as digital—it can be measured, quantified and programmed into the model. Such biological information starts with an organism's genetic code. Every cell in the human body carries a full copy of the human genome, which is made up of three billion pairs of DNA bases, the letters of the genetic alphabet. Those “letters” encode some 25,000 genes, representing instructions for operating cells and tissues. Inside each cell, genes are transcribed into a more portable form, discrete snippets of messenger RNA that carry those instructions to cellular equipment that reads the RNA and churns out chains of amino acids according to the encoded instructions. Those amino acid chains, in turn, fold themselves into proteins, the three-dimensional molecular machines that execute most of the functions of life.

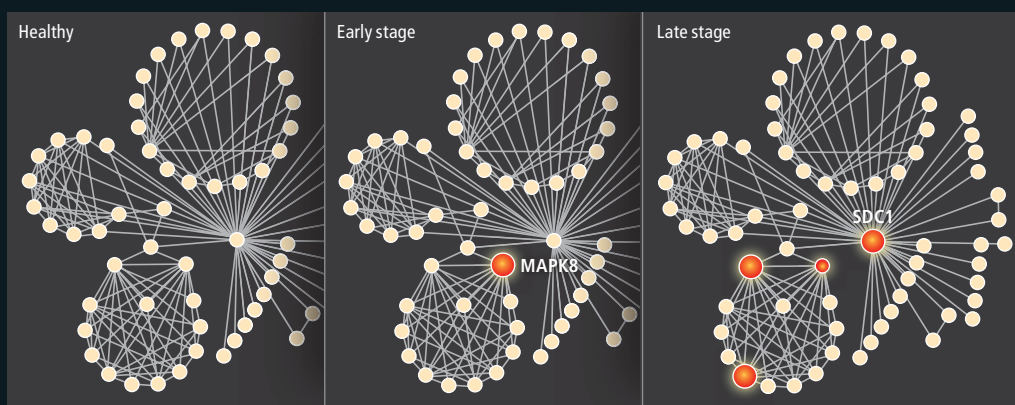
Within a biological system, such as a person, all these “data” are transmitted, processed, integrated and ultimately executed through networks of proteins interacting with one another

and with other biologically relevant molecules inside cells. And when the entire system is viewed as a network of these interrelated events, disease can be seen as a consequence of something perturbing the network's normal programmed patterns of information. The initial cause could be a flaw within the system, such as a random change in DNA that alters an encoded instruction, or even some environmental influence impinging on the system from outside, such as the ultraviolet radiation in sunlight that can trigger DNA damage that eventually leads to melanoma. As an initial disruption produces ripple effects and feedback, the information patterns continue to change, and these dynamically altered patterns explain the nature of the disease mechanistically [see illustration on next page].

Of course, building an accurate computer model of this kind of biological network is a staggering effort. The task can require computational integration of millions or more measurements of messenger RNA and protein levels to comprehensively capture the dynamics of the system's transition from health to disease. Nevertheless, an accurate model—meaning one ca-

NANOPARTICLES CONSTRUCTED to carry a therapeutic payload are studded with proteins that act as keys for gaining entry to tumor cells.

PROSTATE CELLS contain groups of proteins (*solid circles*) that interact (*lines*) with one another in small networks; changes in cellular levels of certain proteins accompany a shift from health to disease. Early-stage prostate cancer cells show a rise in levels of MAPK8, a protein known to regulate cell movement. In late-stage cancer cells, levels of SDC1 are 16 times higher than in early-stage cells. Relative amounts of these two proteins can offer diagnostic clues to the presence and progression of disease.



[THE AUTHORS]

James R. Heath is director of the NanoSystems Biology Cancer Center and professor of chemistry at the California Institute of Technology, where he works on nanostructured materials and nanoelectronic circuitry as well as technologies for cancer diagnosis and treatment. **Mark E. Davis**, a Caltech professor of chemical engineering, develops specialized materials for experimental therapeutics and founded two companies, Insert Therapeutics and Calando Pharmaceuticals, that develop nanoparticle therapies. **Leroy Hood** is president of the Institute for Systems Biology in Seattle, which he founded after having pioneered technologies for DNA and protein sequencing and synthesis and starting numerous companies, including Amgen, Applied Biosystems, Systemix, Darwin and Rosetta. Hood and Heath also founded Integrated Diagnostics, a systems medicine company that is seeking biomarkers for disease and developing microfluidic and nanotechnology platforms for translating those biomarkers into diagnostic tools.

pable of correctly predicting the effects of perturbations—can be the foundation for dramatic changes in the way illness and health are understood and how they are approached medically.

Over the past several decades, for example, cancer has been the most intensively studied of all diseases, yet tumors have traditionally been characterized by fairly coarse features that include their size, their location in a particular organ or tissue, and whether malignant cells have spread from the primary tumor. The more advanced the cancer according to such diagnostic “stages,” the more bleak the prognosis for the patient. But even that conventional wisdom offers plenty of contradictions. Patients diagnosed with identical cancers and given similar treatments from the standard repertoire of radiation and chemotherapies often respond very differently—one group of patients may experience full recovery, whereas a second group succumbs rapidly.

Large-scale measurements of messenger RNA and protein concentrations within biopsied tumors have revealed the inadequacy of such traditional approaches by showing how two patients’ cancers that are seemingly the same actually contain networks perturbed in dramatically different ways. Based on such molecular analysis, many cancers that were at one time considered to be a single disease are now identified as separate diseases.

About 80 percent of human prostate tumors grow so slowly that they will not ever harm their hosts, for instance. The remaining 20 percent will grow more quickly, invading surrounding tissues and even spreading (metastasizing) to distant organs, eventually killing the patient. Our research group is now attempting to identify the disease-perturbed networks in prostate cells that characterize both these major cancer types so that a doctor could identify from the

outset which kind a patient has. That information could spare 80 percent of patients unnecessary surgery, irradiation or chemotherapy, along with the pain, incontinence and impotence that accompany those treatments.

We are also analyzing the networks within the prostate that distinguish subtypes among the more aggressive 20 percent of cases that might require distinct treatment regimens. For example, by analyzing the networks characteristic of early-stage and metastatic prostate cancers, we have identified a protein secreted into the blood that appears to be an excellent identifying marker for metastatic cancer. Tools of this kind that can stratify a given disease such as prostate cancer into its precise subtype would allow a clinician to make a rational choice of the appropriate therapy for each individual.

Detecting Disease

Although such analyses of messenger RNAs and proteins from tumor tissues can be informative about the nature of a known cancer, the systems approach can also be applied to distinguishing between health and illness. Blood bathes every organ in the body, carrying away proteins and other molecules, so it provides an excellent window into the entire body system. The ability to detect an imbalance in particular proteins or messenger RNAs could therefore serve to signal the presence of disease and pinpoint its location and its nature.

Our research group has addressed the challenge of using blood to assess the status of the whole-body system by comparing messenger RNA populations produced in the 50 or so individual organs, and we have found that each human organ has 50 or more messenger RNA types that are made primarily in just that organ. Certain of these RNAs encode organ-specific pro-

JEN CHRISTIANSEN; SOURCE: LEROY HOOD

teins that are secreted into the bloodstream, and the levels of each will reflect the operation of the networks that control their production within the organ. When those networks are disease-perturbed, corresponding protein levels will be altered. Those changes should make it possible to identify the disease because each disease in the organ will disturb distinct biological networks in unique ways.

If the levels of about 25 proteins from each of these organ-specific fingerprints can be assessed, computational analysis should make it possible to detect all diseases by determining which networks are perturbed—just from blood measurements. Beyond early detection—which is so important in cancer—this approach would provide the ability to stratify a patient's disease into its different subtypes, to follow its progression and to follow its response to therapy. We have shown an initial proof of this principle by following the evolution of prion disease in mice.

We injected the mice with infectious prion proteins that lead to a degenerative brain disease akin to “mad cow disease,” then we analyzed the complete populations of brain messenger RNAs

in the infected and control animals at 10 different time points during the onset of the disease. From these data we identified 300 changing messenger RNAs that encoded the core prion disease response. About 200 of those RNAs belonged to four biological networks that explained virtually every known aspect of the disease, and about 100 other RNAs described previously unknown aspects of prion disease. Study of these disease-perturbed networks also allowed us to identify four blood proteins that predicted the presence of prion disease before any apparent symptoms and could therefore serve as presymptomatic diagnostic markers, with obvious benefits for preventive medicine.

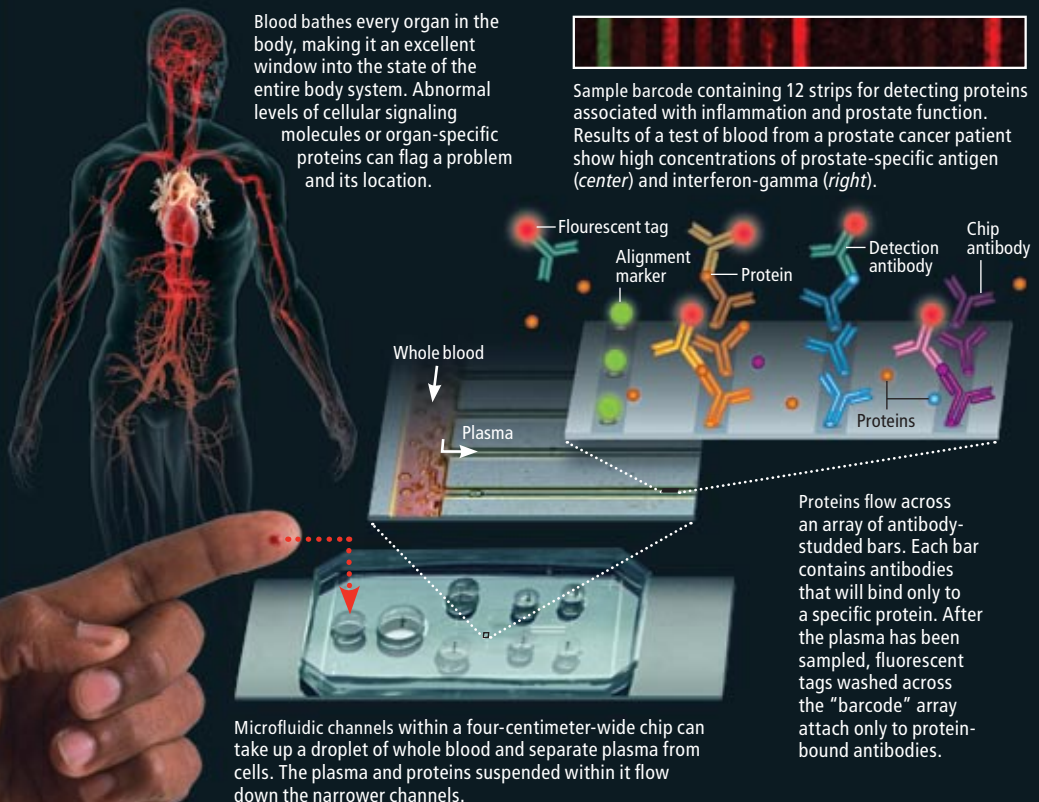
These studies required about 30 million measurements, and we developed a series of software programs for analyzing, integrating and finally modeling these enormous amounts of data. Constructing predictive network models of disease and translating those models into medically useful tools will require rapid, sensitive and—most important—cheap methods for DNA sequencing and for measuring messenger RNA and protein concentrations.

An imbalance in particular proteins or messenger RNAs could serve to signal the presence of disease and pinpoint its location and nature.

[DIAGNOSTICS]

PENNIES PER PROTEIN

Information is the most valuable commodity in a systems approach to medicine, so diagnostic tests will have to easily and accurately measure large numbers of biological molecules for a few cents or less per measurement. Extreme miniaturization allowed the authors and their colleagues to produce a prototype chip that can measure concentrations of a panel of cancer-associated proteins in a droplet of blood in 10 minutes, at a cost of five to 10 cents per protein.

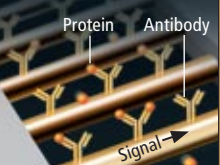
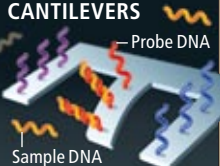
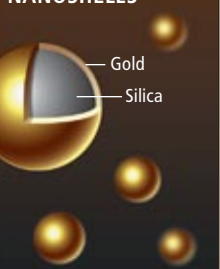
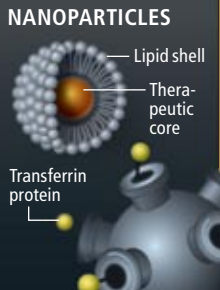


NANOTECH IN MEDICINE

At the scale of one nanometer—one billionth of a meter—materials and devices can interact with cells and biological molecules in unique ways. The nanoscale technologies already used in research or therapies are generally between 10 nanometers, the size of an antibody protein, and 100 nanometers, the size of a virus. These devices and particles are being applied as sensors to detect molecules such as proteins or DNA, as imaging enhancers, and as a means to target specific tissues and deliver therapeutic agents.

◀────────▶
Nanodevices

0.01 nanometer	1	10	100	1,000	10,000	100,000
Glucose		Antibody	Virus	Bacterium	Red blood cell	Hair diameter

NANOTECHNOLOGY	USE	HOW IT WORKS
NANOWIRES 	Sensing	Conductive wire, 10 to 20 nanometers thick, is strung across a channel through which a sample will pass. To detect proteins or DNA, probes made of complementary antibodies or DNA are attached to each wire. When a protein meets its matching antibody, it binds to the probe and changes the conductive properties of the wire, allowing the event to be detected electronically.
CANTILEVERS 	Sensing	Molecular probes, such as single-stranded DNA, can also be attached to beams just a few nanometers thick. When exposed to a DNA sample, complementary strands bind to the probes on the cantilever, causing the beams to bend slightly. That response can be detected visually or by a change in the beams' electrical conductivity.
QUANTUM DOTS 	Imaging	Nanocrystals made of inorganic elements such as cadmium or mercury encased in latex or metal respond to light by emitting fluorescence at different wavelengths and intensities depending on their composition. Antibodies attached to the crystals can cause the dots to bind to a select tissue, such as a tumor, which can then be more easily seen with conventional imaging devices.
NANOSHELLS 	Tissue targeting, Imaging	Solid silica nanospheres, sometimes encased in a thin layer of gold, will travel through the bloodstream without entering most healthy tissues, but they tend to accumulate in tumor tissue. Therapeutic molecules can be attached to the spheres, or once a large number of the nanoshells accumulate in a tumor, heat delivered to the tumor will be absorbed by the spheres, killing the tissue. Depending on their composition, nanoshells can also absorb or scatter light, enhancing tumor images made with certain forms of spectroscopy.
NANOPARTICLES 	Tissue targeting, Delivery	Particles composed of a variety of materials can be constructed to contain therapeutic molecules in their core and to release them at a desirable time and location. Such delivery vehicles include simple lipid shells that passively leak through tumor blood vessel walls, then slowly release a traditional chemotherapy drug into the tissue. Newer nanoparticles are more complexly designed, including exterior elements such as antibodies to target tumor-specific proteins, and materials that minimize the particles' interaction with healthy tissues.

Measuring Molecules

Many scientists have noted that technological advances in DNA sequencing have mirrored Moore's Law for microprocessors: namely, that the number of functional elements that can be placed on a chip per unit cost has doubled every 18 months for the past several decades. In fact, next-generation DNA-sequencing machines are increasing the speed of reading DNA at a pace that is much faster than Moore's Law. For example, the first human genome sequenced probably took about three to four years to complete and cost perhaps \$300 million. We believe within five to 10 years an individual human genome sequence will cost less than \$1,000—a 300,000-fold reduction in price—and be done in a day. Over the next decade, similar advances in other relevant biomedical technologies will allow predictive and personalized medicine to emerge.

At present, a test to measure levels of a single diagnostic cancer protein, such as prostate-specific antigen, in a patient's blood costs a hospital about \$50 to perform. Given that systems-based medicine will require measurements of large numbers of such proteins, this price must fall dramatically. Measurement time is also a cost. A blood test today could take a few hours to a few days, in part because of the many steps needed to separate blood components—cells, plasma, proteins and other molecules—before they can each be measured using tests of varying accuracy.

Extreme miniaturization can offer enhanced precision and significantly faster measurements than can be achieved with current technologies. Several microscale and nanoscale technologies are already proving their value as research tools for gathering the data needed to construct a systems view of biological information. For use in patient care, though, the demands of a systems approach to medicine will mandate that each measurement of a protein should only cost a few pennies—a target that is unlikely to be met by many emerging nanotechnologies.

Two of us (Heath and Hood) have developed a four-centimeter-wide chip that tests protein levels in a droplet of blood, employing a highly miniaturized variant of conventional protein-detection strategies [see box on preceding page]. The chip is made only of glass, plastic and reagents, so it is very inexpensive to produce. Our device takes about two microliters of blood, separates the cells from the plasma, then measures a panel of a dozen plasma proteins, all within a few minutes of blood collection. The projected cost of using the prototype version is perhaps

five to 10 cents per protein tested, but when fully developed, this technology should be able to meet the cost demands of systems medicine.

Extending the capabilities of the chip to measure hundreds of thousands of proteins will take time, but advances in microfluidics design, surface chemistry and measurement science are quickly bridging the gap between what is possible today and what will be required to fully realize a new predictive and personalized medicine. Our Caltech colleagues Stephen R. Quake and Axel Scherer, for instance, have developed a microfluidic system that directly integrates valves and pumps on a chip. Their miniaturized plumbing permits chemical reagents, biomolecules and biological samples to be precisely directed into any one of a large number of individual chambers on the chip, with each chamber representing a separate and independent measurement. Their concept thus turns a lab-on-a-chip into many labs-on-a-chip, providing avenues for further reducing the costs of biological measurements.

Extremely miniaturized technology has similarly important implications for therapies and prevention. Insights into diseased networks can ultimately provide new targets for novel therapies that restore network dynamics to normalcy. In the shorter term, the systems view can help to target existing drugs far more effectively by matching the optimal drug combination to each patient. Nanotechnology, moreover, can also radically reduce the amount of each drug that is required to treat a cancer.

Tiny and Targeted

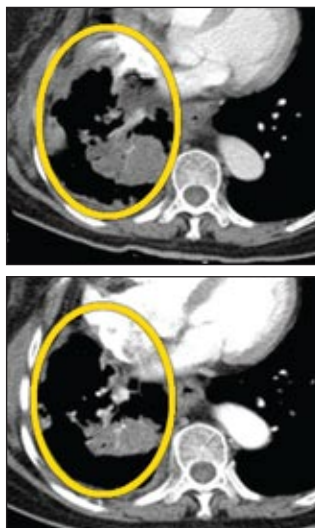
Nanoparticle therapeutics are small relative to most things but large compared to a molecule, and operating at this scale provides an unprecedented level of control over the therapeutic particles' behavior within the body. Nanoparticles can range between one and 100 nanometers (nm) in size and can be assembled from a variety of existing therapeutic agents, such as chemotherapy drugs or gene-silencing RNA (siRNA) strands.

These payloads may be encapsulated in synthetic materials such as polymers or lipidlike molecules, and targeting agents such as antibodies and other molecules designed to bind with specific cellular proteins can be added to a particle's surface. This modularity makes nanotherapeutics especially versatile and capable of performing complex functions at the right place and the right time inside a patient.

Nanotechnology can also radically reduce the amount of each drug required to treat a cancer.

ON TARGET

An experimental nanotherapy, IT-101, encapsulates a chemotherapy drug, camptothecin, inside a nanoparticle designed to circulate for an extended period in the bloodstream and to accumulate in tumors. In a human safety trial, evidence of the treatment's efficacy was seen in some patients with advanced cancers. In the CT scans below, views of a patient's midsection show a large lung tumor (*top*, gray circled mass) before treatment with IT-101 and after six months of treatment (*bottom*), when the tumor had shrunk considerably.



One of the greatest challenges in developing and using cancer drugs is delivering them to the diseased tissues without poisoning the patient's entire body. Size alone gives even simple nanoparticle therapies special properties that determine their movement into and throughout tumors. Nanoparticles smaller than 10 nm are, like so-called small-molecule drugs, rapidly eliminated through the kidney, whereas particles larger than 100 nm have a difficult time moving through a tumor. Particles within the 10- to 100-nm range travel throughout the bloodstream to seek out tumors, although they are unable to escape into most healthy tissues through blood vessel walls. Because tumors, in contrast, have abnormal blood vessels whose walls are riddled with large pores, nanoparticles can leak into the surrounding tumor tissue. As a result, nanoparticles have a tendency to accumulate in tumors while minimizing effects on other parts of the body and avoiding the traditional ravishing side effects of cancer drugs.

Even when a standard drug manages to get to tumor cells, cellular pump proteins may eject it from the cell before it has a chance to work, a common mechanism of drug resistance. Nanoparticles enter a cell by endocytosis, a natural process that creates a pocket of cell membrane around a foreign object to draw it inside the cell, protecting the particle's payload from the cellular pumps [see box on next page].

Certain cancer therapeutics that are now reclassified as nanoparticles have existed for some time and illustrate some of these basic advantages of nanoparticles in reaching tumor cells while minimizing effects on healthy tissues. Liposomal doxorubicin, for example, is a traditional chemotherapy compound encapsulated in a lipid shell that has been used to treat ovarian cancer and multiple myeloma. The lipid-encased version of the drug has far lower heart toxicity than doxorubicin alone, although a new side effect, skin toxicity, has been observed.

Newer nanoparticles—for example one known as IT-101, which has already undergone human safety testing in phase I trials, have more complex designs that provide multiple functions. IT-101 is a 30-nm particle assembled from polymers joined to the small-molecule drug camptothecin, which is closely related to two chemotherapy drugs approved by the Food and Drug Administration: irinotecan and topotecan. The IT-101 particles are designed to circulate in the patient's blood and they remain there for more than 40 hours, whereas camptothecin by itself

would circulate for only a few minutes. This long circulation period allows time for IT-101 to escape into tumors and accumulate there. The particles then enter tumor cells and slowly release the camptothecin to enhance its effects. As the drug is released, the rest of the nanoparticle components disassemble and the small individual polymer molecules harmlessly exit the body through the kidneys.

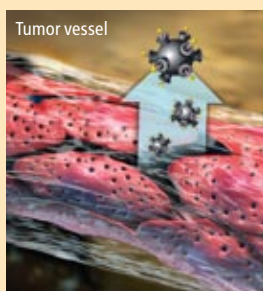
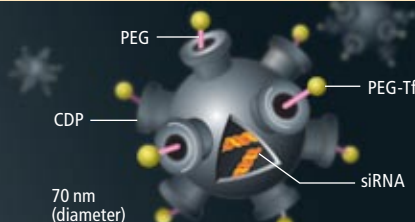
[CASE STUDY]

DESIGNED TO DELIVER

An experimental nanoparticle therapeutic called CALAA-01 illustrates some of the advantages such agents can offer. In addition to having a natural tendency to accumulate in tumors, nanoparticles can be designed to home to one or more receptors commonly found on cancer cells. The particles' mode of cell entry also allows them to evade cellular pumps that eject some drugs.

CUSTOMIZED STRUCTURE

The particle is built with biocompatible materials: a cyclodextrin-containing polymer (CDP) with polyethylene glycol (PEG) stalks to which transferrin proteins (Tf) are attached. Inside, as many as 2,000 siRNA molecules—the therapeutic agents—are stored.

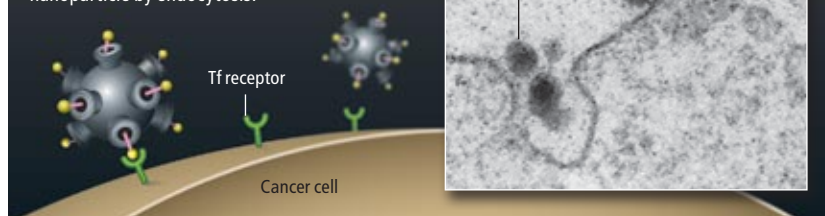


PASSIVE TUMOR TARGETING

When the particles enter a patient's bloodstream, they circulate freely but cannot penetrate most blood vessel walls. Tumor vessels are abnormally leaky, with large pores that allow nanoparticles to pass through and accumulate in the tumor tissue.

ACTIVE TUMOR TARGETING

Transferrin receptors on the surface of a cancer cell bind to the transferrin protein on the nanoparticle, causing the cell to internalize the nanoparticle by endocytosis.



CONTROLLED RELEASE

Once inside the cell, a chemical sensor within the nanoparticle responds to the low pH within the endocytotic vesicle by simultaneously triggering disassembly of the nanoparticle and release of the siRNA molecules that will block a gene's instructions from being translated into a protein the cancer cell needs to survive.



In the clinical trials, dosages of the drug were achieved that provided high quality of life without the side effects, such as vomiting, diarrhea and hair loss typical of chemotherapeutics, and with no new side effects. The general high quality of life while on treatment is exciting, and although phase I trials focus on establishing safety, the tests also provided evidence that the drug was active in patients [see box on preceding page]. That is encouraging because patients in phase I cancer trials have had numerous courses of standard therapy that failed before entering the trial. After completing the six-month trial, several of these patients have remained on the drug on a compassionate use basis, and long-term survivors of approximately one year or more include patients with advanced lung, renal and pancreatic cancers.

Because the side-effect profile of this drug is so low, it will next be tested in a phase II (efficacy) trial in women diagnosed with ovarian cancer who have undergone chemotherapy. Instead of simply “waiting and watching” for the cancer to progress, IT-101 will be given as maintenance therapy in the hope of preventing disease progression. These observations from IT-101 testing and similarly encouraging news from trials of other nanoparticle-based treatments are beginning to provide a picture of what may be possible with well-designed nanotherapeutics. Indeed, the next generation of synthetic nanoparticles, which are far more sophisticated, offers a glimpse of the real potential of this technology and the importance these drugs will hold for the systems-based view of disease and treatment.

Calando Pharmaceuticals in Pasadena, Calif., began trials in 2008 of a siRNA delivery system invented by one of us (Davis) that illustrates the newer approach. Proteins on the surface of the particles target specific receptors that occur in high concentrations on the surface of cancer cells. Once inside the cells, the particles release siRNA molecules, which are tailored to match a specific gene of interest and inhibit the manufacture of the gene's encoded protein.

This multifunctional nanotherapeutic is just the beginning of the story, though. Once the principles of nanoparticle function in humans are more fully established, the concept can be applied to create a therapeutic system that can carry combinations of drugs, each with its own customized release rates. For example, if one wished to inhibit a protein that causes a drug to be ineffective, then an option would be to create

a nanoparticle that first releases siRNA to inhibit the gene for that protein before releasing the drug molecule. As greater understanding of the molecular transitions from health to disease and vice versa is gained, the nanoparticle approach is likely to play an increasing role in the treatment of disease at the molecular level.

The Big Picture

A systems approach to disease is predicated on the idea that the analysis of dynamic disease-perturbed networks and the detailed mechanistic understanding of disease that it provides can transform every aspect of how we practice medicine—better diagnostics, effective new approaches to therapy and even prevention. This systems biology approach to disease is driving the development of many new technologies, including microfluidics, nanotechnologies, new measurement and visualization instrumentation, and computational advances that can analyze, integrate and model large amounts of biological information.

In the next 10 to 20 years, predictive and personalized medicine will be revolutionized by at least two new approaches. The sequence of individual human genomes will permit us to determine with ever increasing accuracy the probable future health of an individual. Inexpensive measurements of blood proteins will permit us to assess, regularly and comprehensively, how that individual's health is evolving.

Preventive medicine starts with the identification of proteins within a diseased network that, if perturbed, will restore network behavior to normalcy, and will eventually lead to prophylactic drugs that prevent disease. For instance, a woman at increased risk for ovarian cancer, who at age 30 starts taking a nanotherapeutic that is specially designed to offset the molecular source of the risk, might lower her lifetime chance of developing ovarian cancer from 40 to 2 percent.

With this kind of knowledge about the causes of health and disease available, people will also be able to participate more effectively in their own health decisions, much as today's diabetics have tools and information that help them manage their own daily well-being.

The realization of a form of medicine that is predictive, personalized, preventive and participatory will have broad implications for society. The health care industry will have to fundamentally alter its business plans, which are currently failing to produce cost effective and highly effi-

NANOSCALE THERAPIES

Nanoscale particles designed to treat cancer include drugs already in use, such as a liposome-encased version of the chemotherapy doxorubicin, as well as a variety of experimental polymer drug combinations in which the drug and polymer molecules are meshed or chemically joined into nanoparticles (composites, conjugates, micelles, dendrimers). Newer targeted particles bear additional features that increase their affinity for, and facilitate entry into, cancer cells.

PARTICLE TYPE	DEVELOPMENT STAGE	EXAMPLES
Liposome	FDA-approved	DaunoXome, Doxil
Albumin-based	FDA-approved	Abraxane
Polymeric micelle	Clinical trials	Genexol-PM, SP1049C, NK911, NK012, NK105, NC-6004
Polymer-drug conjugate	Clinical trials	XYOTAX (CT-2103), CT-2106, IT-101, AP5280, AP5346, FCE28068 (PK1), FCE28069 (PK2), PNU166148, PNU166945, MAG-CPT, DE-310, Pegamotecan, NKTR-102, EZN-2208
Targeted liposome	Clinical trials	MCC-465, MBP-426, SGT-53
Targeted polymer-based particle	Clinical trials	FCE28069(PK2), CALAA-01
Solid inorganic or metal particle	Clinical trials (gold) and preclinical	Carbon nanotubes, silica particles, gold particles (CYT-6091)
Dendrimer	Preclinical	Polyamidoamine (PAMAM)

cacious drugs. Emerging technologies will also lead to the digitization of medicine—meaning the ability to extract disease-relevant information from single molecules, single cells or single individuals—just as information technologies and communications have been digitized over the past 15 years. As a result of the new high-throughput, low-cost technologies, the cost of health care should fall sharply, making it widely accessible even in the developing world.

For cancer, the very exciting promises that should be realized over the next 10 years are, first, that presymptomatic blood diagnoses will catch fledgling cancers that can be cured by conventional therapy; second, that cancers of a particular tissue (for example, breast or prostate) will be stratified into distinct types that can be matched with drugs that provide very high cure rates; and third, the identification of disease-perturbed networks will allow more rapid development of drugs that are less expensive and far more effective. This new approach to medicine therefore has the potential to transform health care for virtually everyone living today. ■

MORE TO EXPLORE

NanoSystems Biology. James R. Heath et al. in *Molecular Imaging and Biology*, Vol. 5, No. 5, pages 312–325; September/October 2003.

Nanotechnology and Cancer. James R. Heath and Mark E. Davis in *Annual Review of Medicine*, Vol. 59, pages 251–265; February 2008. (First published online: October 15, 2007.)

Nanoparticle Therapeutics: An Emerging Treatment Modality for Cancer. Mark E. Davis et al. in *Nature Reviews Drug Discovery*, Vol. 7, No. 9, pages 771–782; September 2008.

Integrated Barcode Chips for Rapid, Multiplexed Analysis of Proteins in Microliter Quantities of Blood. Rong Fan et al. in *Nature Biotechnology*. Advance online publication: November 16, 2008.

The ORIGIN of the Land UNDER the Sea

The deep basins under the oceans are carpeted with lava that spewed from submarine volcanoes and solidified. Scientists have solved the mystery of how, precisely, all that lava reaches the seafloor

By Peter B. Kelemen

KEY CONCEPTS

- Eighty-five percent of the earth's volcanic eruptions occur deep underwater along mid-ocean ridges.
- Lava ejected from those narrow chains of seafloor volcanoes produce the rocky underpinnings of all oceans.
- Until recently, no one understood much about how the molten lava rises up into the ridges.
- Scientists now think they have deciphered the process, beginning with the formation of microscopic droplets of liquid rock in regions up to 150 kilometers deep.

—The Editors

At the dark bottom of our cool oceans, 85 percent of the earth's volcanic eruptions proceed virtually unnoticed. Though unseen, they are hardly insignificant. Submarine volcanoes generate the solid underpinnings of all the world's oceans—massive slabs of rock seven kilometers thick.

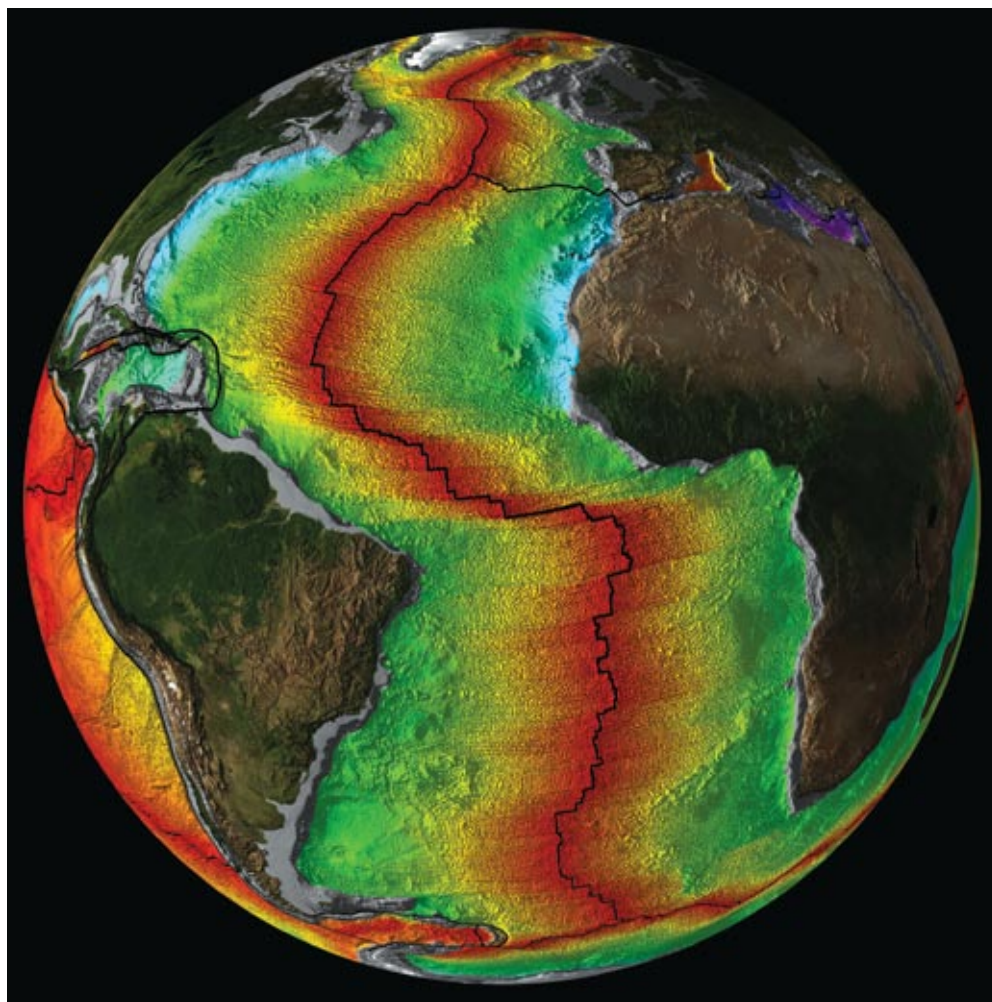
Geophysicists first began to appreciate the smoldering origins of the land under the sea, known formally as ocean crust, in the early 1960s. Sonar surveys revealed that volcanoes form nearly continuous ridges that wind around the globe like seams on a baseball. Later, the same scientists strove to explain what fuels these erupting mountain ranges, called mid-ocean ridges. Basic theories suggest that because ocean crust pulls apart along the ridges, hot material deep within the earth's rocky interior must rise to fill the gap. But details of exactly where the lava originates and how it travels to the surface long remained a mystery.

In recent years mathematical models of the interaction between molten and solid rock have provided some answers, as have examinations of blocks of old seafloor now exposed on the continents. These insights made it possible to devel-

op a detailed theory describing the birth of ocean crust. The process turns out to be quite different from the typical layperson's idea, in which fiery magma fills an enormous chamber underneath a volcano, then rages upward along a jagged crack. Instead the process begins dozens of kilometers under the seafloor, where tiny droplets of melted rock ooze through microscopic pores at a rate of about 10 centimeters a year, about as fast as fingernails grow. Closer to the surface, the process speeds up, culminating with massive streams of lava pouring over the seafloor with the velocity of a speeding truck. Deciphering how liquid moves through solid rock deep underground not only explains how ocean crust emerges but also may elucidate the behavior of other fluid-transport networks, including the river systems that dissect the planet's surface.

Digging Deep

Far below the mid-ocean ridge volcanoes and their countless layers of crust-forming lava is the mantle, a 3,200-kilometer-thick layer of scorching hot rock that forms the earth's midsection and surrounds its metallic core. At the planet's cool surface, upthrust mantle rocks are dark



MID-ATLANTIC RIDGE (black "stitching"), a 10,000-kilometer-long string of volcanoes at the bottom of the Atlantic Ocean, is the longest mountain range in the world. The colors indicate the ages of the rocky crust under the ocean, which is youngest (red) near the ridge and gets progressively older as it nears the continents.

GEOLOGIC GLOSSARY

DUNITE: A rock that is made up almost entirely of the mineral olivine and typically forms a network of light-colored veins in the upper mantle.

FOCUSED POROUS FLOW: The process by which melt often travels through the solid rocks of the earth's interior; the melt flows through elongated pores between single microscopic crystals, somewhat like water flowing through sand.

LAVA: Molten rock expelled during a volcanic eruption at the earth's surface.

MELT: The name for molten rock before it erupts.

MID-OCEAN RIDGES: Underwater mountain ranges that generate new ocean crust through volcanic eruptions.

MINERALS: The building blocks of rocks, which can be made up of single elements, such as gold, or multiple elements; olivine, for instance, comprises magnesium, silicon and oxygen.

OPHOLITE: A section of ocean crust and underlying mantle rocks that has been uplifted onto the continents during the collision of tectonic plates.

green, but if you could see them in their rightful home, they would be glowing red- or even white-hot. The top of the mantle is about 1,300 degrees Celsius, and it gets about one degree hotter with each kilometer of depth. The weight of overlying rock means the pressure also increases with depth—about 1,000 atmospheres for every three kilometers.

Knowledge of the intense heat and pressure in the mantle led researchers to hypothesize in the late 1960s that ocean crust originates as tiny amounts of liquid rock known as melt—almost as though the solid rocks were "sweating." Even a minuscule release of pressure (because of material rising from its original position) causes melt to form in microscopic pores deep within the mantle rock.

Explaining how the rock sweat gets to the surface was more difficult. Melt is less dense than the mantle rocks in which it forms, so it will constantly try to migrate upward, toward regions of lower pressure. But what laboratory experiments revealed about the chemical composition of melt did not seem to match up with the composition of rock samples collected from the mid-ocean ridges, where erupted melt hardens.

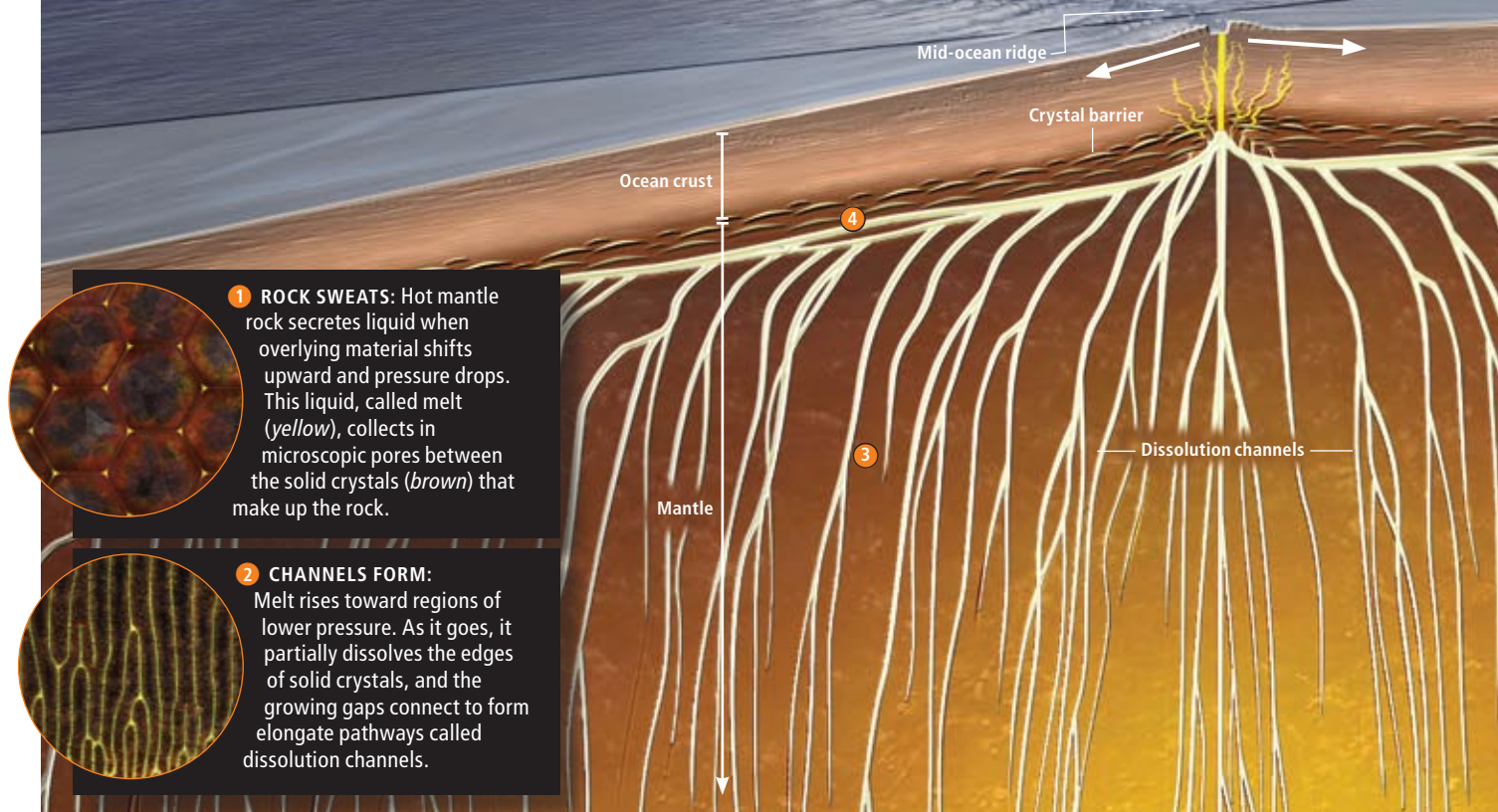
Using specialized equipment to heat and squeeze crystals from mantle rocks in the laboratory, investigators learned that the chemical composition of melt in the mantle varies depending on the depth at which it forms; the composition is controlled by an exchange of atoms between the melt and the minerals that make up the solid rock it passes through. The experiments revealed that as melt rises, it dissolves one kind of mineral, orthopyroxene, and precipitates, or leaves behind, another mineral, olivine. Researchers could thus infer that the higher in the mantle melt formed, the more orthopyroxene it would dissolve, and the more olivine it would leave behind. Comparing these experimental findings with lava samples from the mid-ocean ridges revealed that almost all of them have the composition of melts that formed at depths greater than 45 kilometers.

This conclusion spurred a lively debate about how melt is able to rise through tens of kilometers of overlying rock while preserving the composition appropriate for a greater depth. If melt rose slowly in small pores in the rock, as researchers suspected, it would be logical to assume that all melts would reflect the composition of the

SWEATING THE SEAFLOOR

The solid foundation of the earth's ocean basins consists of massive slabs of volcanic rock seven kilometers thick. This rock, known as ocean crust, originates as tiny droplets that form across broad regions of the planet's solid interior, or man-

tle, almost as though the rocks were sweating. Yet through a process known as focused porous flow, those countless droplets all emerge along mid-ocean ridges. As this new material rises up, older crust moves away from the ridges (arrows).



1 ROCK SWEATS: Hot mantle rock secretes liquid when overlying material shifts upward and pressure drops. This liquid, called melt (yellow), collects in microscopic pores between the solid crystals (brown) that make up the rock.

2 CHANNELS FORM: Melt rises toward regions of lower pressure. As it goes, it partially dissolves the edges of solid crystals, and the growing gaps connect to form elongate pathways called dissolution channels.

3 MELT OOZES SLOWLY: Melt rises only a few centimeters a year because dissolution channels are congested with rock grains that the melt cannot dissolve. Gradually, millions of threads of melt coalesce into larger conduits.

4 BARRIERS BLOCK FLOW: In the cool upper mantle, some rising melt loses enough heat to crystallize, forming solid barriers. These crystal barriers arise deeper far away from the ridge, so they guide the remaining melt toward the ridge.

5 CRACKS OPEN: Underneath the mid-ocean ridge, crystallized melt blocks upward flow completely. Melt accumulates in lens-shaped pockets until the pressure inside rises enough to fracture the colder rocks above.

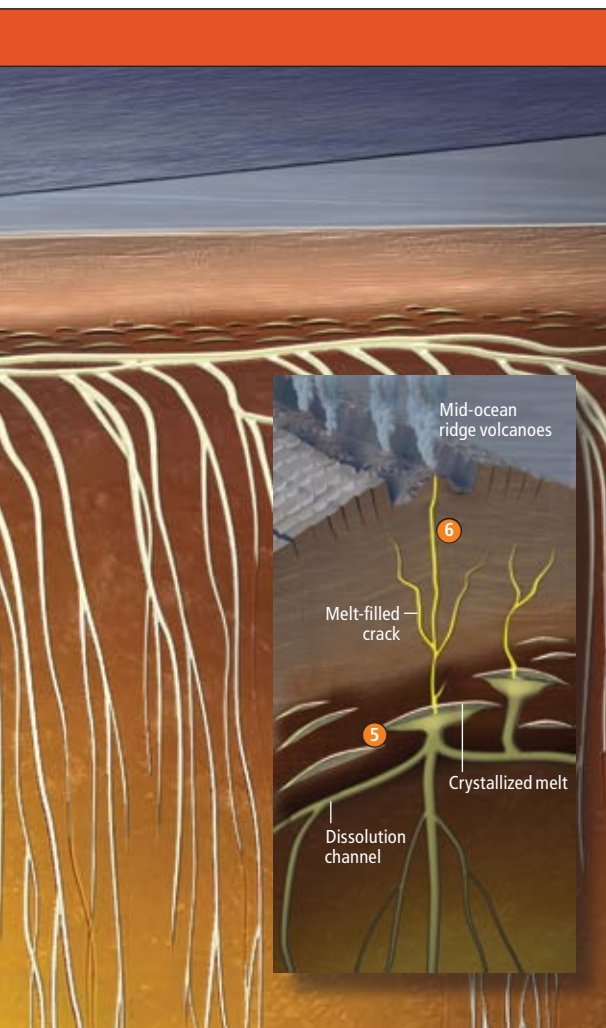
shallowest part of the mantle, at 10 kilometers or less. Yet the composition of most mid-ocean ridge lava samples suggests their source melt migrated through the uppermost 45 kilometers of the mantle without dissolving any orthopyroxene from the surrounding rock. But how?

Cracking under Pressure?

In the early 1970s scientists proposed an answer that was not unlike the typical layperson's view: the melt must make the last leg of its upward journey along enormous cracks. Open cracks would allow the melt to rise so rapidly that it would not have time to interact with the surrounding rock, nor would melt in the core of the crack ever touch the sides. Although open cracks are not a natural feature of the upper mantle—

the pressure is simply too great—some investigators suggested that the buoyant force of migrating melt might sometimes be enough to fracture the solid rock above, like an icebreaker ship forcing its way through polar pack ice.

Adolphe Nicolas of the University of Montpellier in France and his colleagues discovered tantalizing evidence for such cracks while examining unusual rock formations called ophiolites. Typically, when oceanic crust gets old and cold, it becomes so dense that it sinks back into the mantle along deep trenches known as subduction zones, such as those that encircle the Pacific Ocean. Ophiolites, on the other hand, are thick sections of old seafloor and adjacent, underlying mantle that are thrust up onto continents when two of the planet's tectonic plates collide. A fa-



6 MELT FLOWS QUICKLY: Open cracks allow melt to escape rapidly upward, completely draining the pocket below. Some melt erupts as lava atop the mid-ocean ridge volcanoes, but most of it crystallizes within the crust below.

amous example, located in the Sultanate of Oman, was exposed during the ongoing collision of the Arabian and Eurasian plates. In this and other ophiolites, Nicolas's team found unusual, light-colored veins called dikes, which they interpreted as cracks in which melt had crystallized before reaching the seafloor.

The problem with this interpretation was that the dikes are filled with rock that crystallized from a melt that formed in the uppermost reaches of the mantle, not below 45 kilometers, where most mid-ocean ridge lavas originate. In addition, the icebreaker scenario may not work well for the melting region under mid-ocean ridges: below about 10 kilometers, the hot mantle tends to flow like caramel left too long in the sun, rather than cracking easily.

A Porous River

To explain the ongoing mystery, I began working on an alternative hypothesis for lava transport in the melting region. In my dissertation in the late 1980s, I developed a chemical theory proposing that as rising melt dissolves orthopyroxene crystals, it precipitates a smaller amount of olivine, so that the net result is a greater volume of melt. In the 1990s my colleagues—Jack Whitehead of the Woods Hole Oceanographic Institution; Einat Aharonov, now at the Weizmann Institute of Science in Rehovot, Israel; and Marc Spiegelman of Columbia University's Lamont-Doherty Earth Observatory—and I created a mathematical model of this process. Our calculations revealed how this dissolution process gradually enlarges the open spaces at the edges of solid crystals, creating larger pores and carving a more favorable pathway through which melt can flow.

As the pores grow, they connect to form elongate channels. In turn, similar feedbacks drive the coalescence of several small tributaries to form larger channels. Indeed, our numerical models suggested that more than 90 percent of the melt is concentrated into less than 10 percent of the available area. That means millions of microscopic threads of flowing melt may eventually feed into only a few dozen, high-porosity channels 100 meters or more wide.

Even in the widest channels, many crystals of the original mantle rock remain intact, congesting the channels and inhibiting movement of the fluid. That is why melt flows slowly, at only a few centimeters a year. Over time, however, so much melt passes through the channels that all the soluble orthopyroxene crystals dissolve away, leaving only crystals of olivine and other minerals that the melt is unable to dissolve. As a result, the composition of the melt within such channels can no longer adjust to decreasing pressure and instead records the depth at which it last "saw" an orthopyroxene crystal.

One of the most important implications of this process, called focused porous flow, is that only the melt at the edges of channels dissolves orthopyroxene from the surrounding rock; melt within the inner part of the conduit can rise unadulterated. Numerical models thus provided key evidence that melt forming deep within the mantle could forge its own upward path—not by cracking the rock but by dissolving some of it. These modeling results were complemented by our ongoing fieldwork, which yielded more direct evidence for porous flow in ophiolites.

FAST FACTS

- It takes an average of 100 years to generate a new swath of oceanic crust six meters wide and seven kilometers thick.
- The hot rock of the earth's mantle is composed of solid crystals, but like a glacier made of solid crystals of ice, mantle rock can flow up to 10 centimeters a year—about as fast as fingernails grow.
- Eruptions of lava at mid-ocean ridges billow across the seafloor at velocities that may at times surpass 100 kilometers per hour.

HOW MID-OCEAN RIDGES FORM

Sometimes new oceans are born on dry land, where tectonic forces pull the land apart in a process called continental rifting. A rising plume of hot mantle rock breaks up and thins the continent from below. This thinning segment of continental crust constitutes a so-called rift valley in which eruptions of basaltic lava—the same kind of lava that makes up ocean crust—become common. As the two sides of the rift diverge, the valley eventually falls below sea level and floods. Once underwater, the volcanoes of the former rift valley become a mid-ocean ridge, where new ocean crust is generated between the two pieces of the old continent.

[INTRIGUING SIMILARITIES]

Water on Land Flows Like Melt in the Mantle

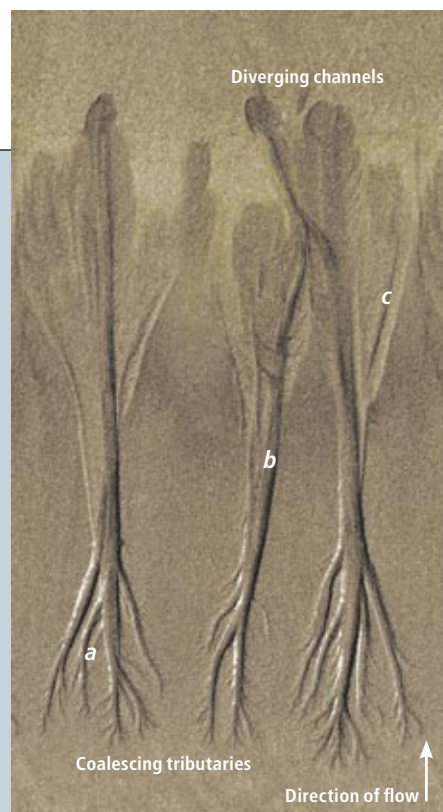
Water flowing over a beach carves a network of channels, much like those that molten rock, or melt, forms as it rises through the earth's solid interior. Although these channel patterns develop by different means—water on the beach physically lifts sand grains and moves them, whereas melt dissolves some of the surrounding rock—their likeness suggests that similar physical laws govern each.

In both cases, regular patterns arise from random starting conditions. On the beach, groundwater that emerges on the surface at low tide flows quickly toward any low spots. This water picks up sand grains, and as it flows, it carves increasingly deep channels that in turn redirect any additional water that enters them along the way (a). As a result of this feed-back mechanism, regularly spaced channels arise and coalesce downstream (b)—even though their point of origin is random. Like creeks and streams feeding a large river, this erosion pattern emerges because it is the best way for the system to conserve energy: the deeper and wider a channel is, the less energy is lost because of friction between the moving water and underlying sand.

Conservation of energy is also a key factor controlling chemical erosion in the earth's mantle. As melt dissolves the surrounding rock, it gradually enlarges the tiny open spaces, or pores, it travels through. These so-called dissolution channels grow and coalesce as the melt migrates higher because the viscous drag, which is analogous to friction, diminishes as pores grow larger [see box on preceding two pages]. As on the beach, many small, active channels supply a few large ones—a pattern that partially explains why undersea eruptions are generally confined to mid-ocean ridges rather than scattered randomly across the seafloor.

Changing conditions can force coalescing channels to diverge. On the beach, water drops its load of sand as the slope flattens out, forming barriers that divert water away from the main channel (c). As in a delta where a river meets the sea, water accumulates behind the obstructions, then periodically overflows and erodes new pathways, which in turn become clogged and abandoned. A similar divergence occurs in the uppermost mantle, which is cooler than the rock deeper down. There the melt cannot remain entirely molten, and parts of it crystallize. Periodically, though, melt bursts through the solid crystal barriers—sometimes forging a new conduit to the seafloor.

—P.B.K.



Focused Effort

The only way to fully appreciate the Oman ophiolite is from the air. This massive formation constitutes a nearly continuous band of rock 500 kilometers long and up to 100 kilometers wide. Like all ophiolites, the mantle part of the Oman ophiolite is generally weathered to a rusty brown color and conspicuously laced with thousands of veins of tan-colored rock. Geologists had long ago identified these veins as a rock called dunite but had not carefully measured the compositions of the minerals within either the dunite or the surrounding rock.

As scientists would expect for rocks once part of the upper mantle, the surrounding rock is rich in both olivine and orthopyroxene. The dunite, on the other hand, is more than 95 percent olivine—the mineral left behind as melt rises through the mantle. Dunite also completely lacks orthopyroxene, which is consistent with the chemical theory predicting that all the orthopyroxene would have been dissolved away by the time the melt reached the uppermost mantle. From this and other evidence, it seemed clear that the dunite veins were conduits that

transported deep melts up through the shallow mantle beneath a mid-ocean ridge. We were seeing dissolution channels frozen in time.

As exciting as these discoveries were, they did not fully explain a second mystery that long perplexed geophysicists. The massive lava flows at mid-ocean ridges emerge from a zone only about five kilometers wide. Yet seismic surveys, which can distinguish between solid and partially molten rock, show that melt exists to a depth of at least 100 kilometers in a region several hundred kilometers across. How, then, does rising lava get channeled into such a narrow region of volcanism on the seafloor?

In 1991 David Sparks and Mark Parmentier, both then at Brown University, proposed an answer rooted in the variable temperature of ocean crust and the uppermost mantle. Newly erupting lava constantly adds material to the slabs of ocean crust on either side of a mid-ocean ridge. As older parts of the slabs move away from the ridge, making way for new, hot lava, they gradually cool. The cooler the crust, the denser it becomes, and so the farther into the warm mantle it sinks. This cooling trend means that in the

ON THE WEB

Witness the birth of a new ocean in photographs and illustrations at www.SciAm.com/ocean-photos

PETER B. KELEVEN

open ocean far from the crest of a mid-ocean ridge, the seafloor and base of the ocean crust below it are an average of about two kilometers deeper than the seafloor and the base of the crust directly underneath the ridge. In addition, the cold crust cools the top of the mantle, so that the cooled part of the uppermost mantle gets thicker, and its base deeper, farther from the ridge.

Based on this relation, Sparks and Parmentier created a computer model of porous flow within the mantle. In their simulations, they could see that some of the rising melt loses enough heat to crystallize in the uppermost mantle—effectively forming a dam or roof. These barriers form deeper the farther from the hot, mid-ocean ridges they are, so as the remaining melt migrates upward, it is forced to do so at an angle, following this inclined roof toward the ridge.

Ultimate Eruptions

Field observations and theoretical models thus provided good explanations for two major mysteries. Ascending melt does not take on the chemical composition in equilibrium with the surrounding mantle rocks because it is chemically isolated within wide dunite conduits. And these conduits are directed toward the mid-ocean ridges as some of the melt cools and crystallizes in the uppermost mantle. But a new question soon arose: If the rise of melt is a continuous, gradual process, as we predict, what unleashes the periodic bursts of molten rock that constitute volcanic eruptions on the seafloor?

Again, field geology guided our theorizing. In the Oman ophiolite, Nicolas and his colleague

Françoise Boudier at Montpellier showed in the mid-1990s that melt accumulates in lens-shaped pockets—a few meters to tens of meters high and tens to hundreds of meters wide—within the shallowest part of the mantle, just below the base of the ocean crust. To explain the physical processes involved, my colleagues and I had to consider how differently mantle rocks behave just below the base of the crust than they do at greater depth.

Underneath an active spreading ridge such as the East Pacific Rise or the ridge that formed the Oman ophiolite, the rocks of the uppermost mantle (that is, the part of the mantle within 2,000 meters of the base of the crust) lose heat to the cool, overlying seafloor. Consequently, some melt cools and crystallizes. With the constant influx of melt from below blocked from continuing upward, melt begins accumulating in lens-shaped pockets below the crystallized material. As more melt enters, the pressure inside the lens rises. Deeper down, the rocks would be hot enough to flow in response, thereby relieving this pressure surge, but here heat loss to the overlying seafloor makes the rocks too stiff. As a result of the increasing pressure, the rocks above the melt pockets periodically crack, forming conduits that lead into the young ocean crust above. Some of the melt collects and freezes near the base of the crust, building up new rocks without ever erupting. But at times the melt surges all the way up and out the neck of a volcano, forming a lava flow up to 10 meters thick and 10 kilometers long, incrementally paving the seafloor with volcanic rock.

Branching Out

These detailed insights into melt-transport networks deep underneath the seafloor bear many similarities to what scientists know about river networks on the earth's surface.

Like the force of small streams joining to form rivers, chemical erosion in the deeper mantle forms a coalescing network in which many small, active tributaries feed fewer, larger channels. Melt that crystallizes in the upper mantle forms “levees” that redirect its flow much like a muddy river that deposits sediment and creates natural levees as it flows into the sea. In both situations, the levees are breached periodically, allowing large, transient outbursts to pass through a single conduit. Research on the underlying physical processes that govern river- and melt-transport networks may ultimately yield a single fundamental theory that explains the behavior of both. ■

[THE AUTHOR]

Peter B. Kelemen is Arthur D. Storke Memorial Professor at Columbia University's Lamont-Doherty Earth Observatory. Kelemen first began contemplating how magma migrates through the earth's rocky interior in 1980 while he was in the Indian Himalaya mapping ophiolites. Since then, he has investigated the interaction between solid and molten rock using a combination of fieldwork, mathematical chemical modeling and fluid dynamics research.

MORE TO EXPLORE

Extraction of Mid-Ocean Ridge Basalt from Upwelling Mantle by Focused Flow of Melt in Dunite Channels. Peter B. Kelemen, Nobumichi Shimizu and Vincent J. M. Salters in *Nature*, Vol. 375, pages 747–753; June 29, 1995.

Causes and Consequences of Flow Organization during Melt Transport: The Reaction Infiltration Instability. Marc W. Spiegelman, Peter B. Kelemen and Einat Aharonov in *Journal of Geophysical Research*, Vol. 106, No. B2, pages 2061–2077; 2001.

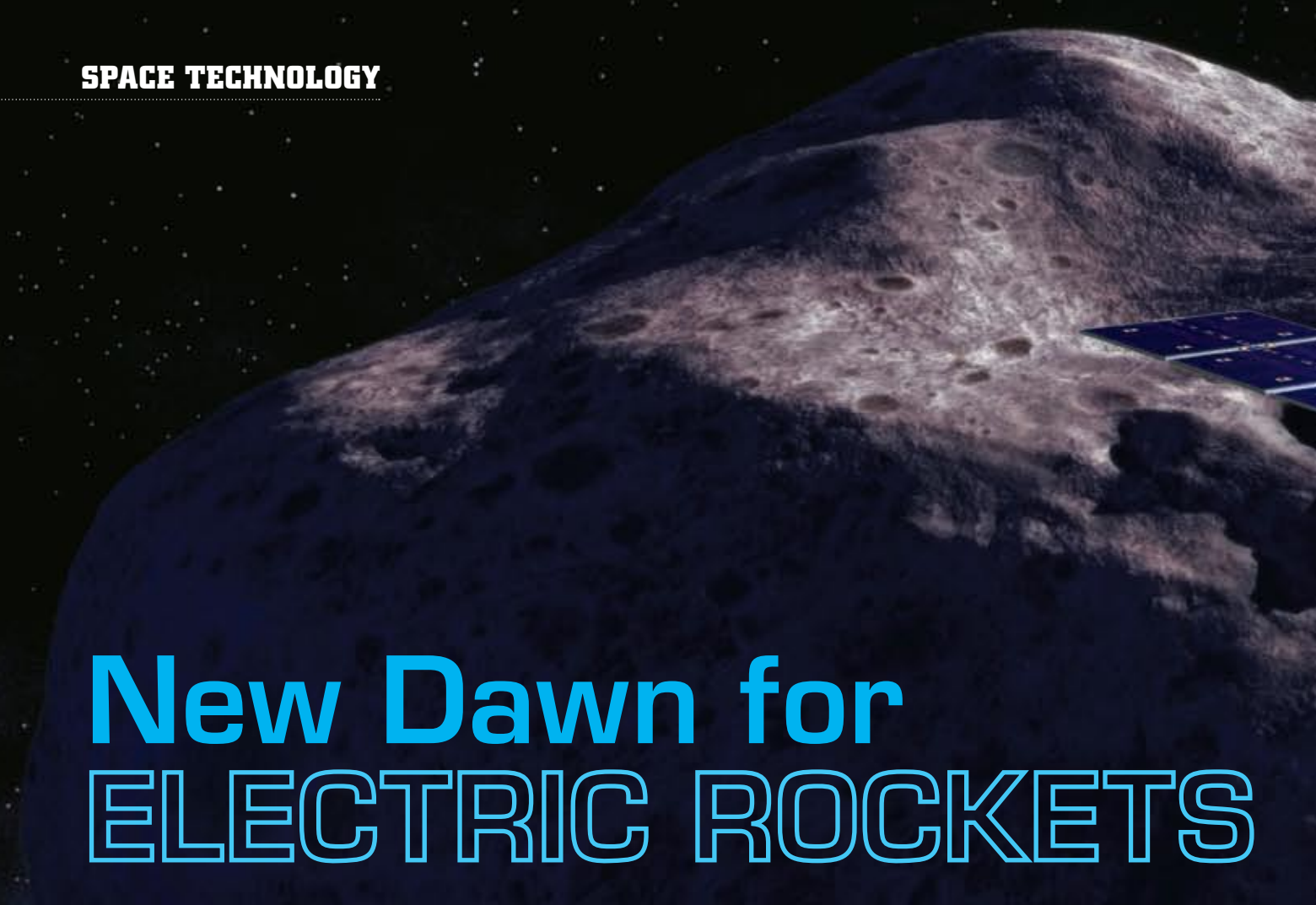
Dunite Distribution in the Oman Ophiolite: Implications for Melt Flux through Porous Dunite Conduits. Michael G. Braun and Peter B. Kelemen in *Geochemistry, Geophysics, Geosystems*, Vol. 3, No. 11; November 6, 2002.

See movies of mathematical models of seafloor spreading at www.ideo.columbia.edu/~mspieg/ReactiveFlow



FORMER OCEAN CRUST and part of the underlying mantle are exposed on dry land in the Sultanate of Oman. Called an ophiolite, this massive brown rock formation, now weathered into rugged hills, was thrust up onto the continent during the ongoing collision of two tectonic plates.

BRADLEY R. HACKER



New Dawn for ELECTRIC ROCKETS

KEY CONCEPTS

- Conventional rockets generate thrust by burning chemical fuel. Electric rockets propel space vehicles by applying electric or electromagnetic fields to clouds of charged particles, or plasmas, to accelerate them.
- Although electric rockets offer much lower thrust levels than their chemical cousins, they can eventually enable spacecraft to reach greater speeds for the same amount of propellant.
- Electric rockets' high-speed capabilities and their efficient use of propellant make them valuable for deep-space missions.

—The Editors

Efficient electric plasma engines are propelling the next generation of space probes to the outer solar system

By Edgar Y. Choueiri

A lone amid the cosmic blackness, NASA's Dawn space probe speeds beyond the orbit of Mars toward the asteroid belt. Launched to search for insights into the birth of the solar system, the robotic spacecraft is on its way to study the asteroids Vesta and Ceres, two of the largest remnants of the planetary embryos that collided and combined some 4.57 billion years ago to form today's planets.

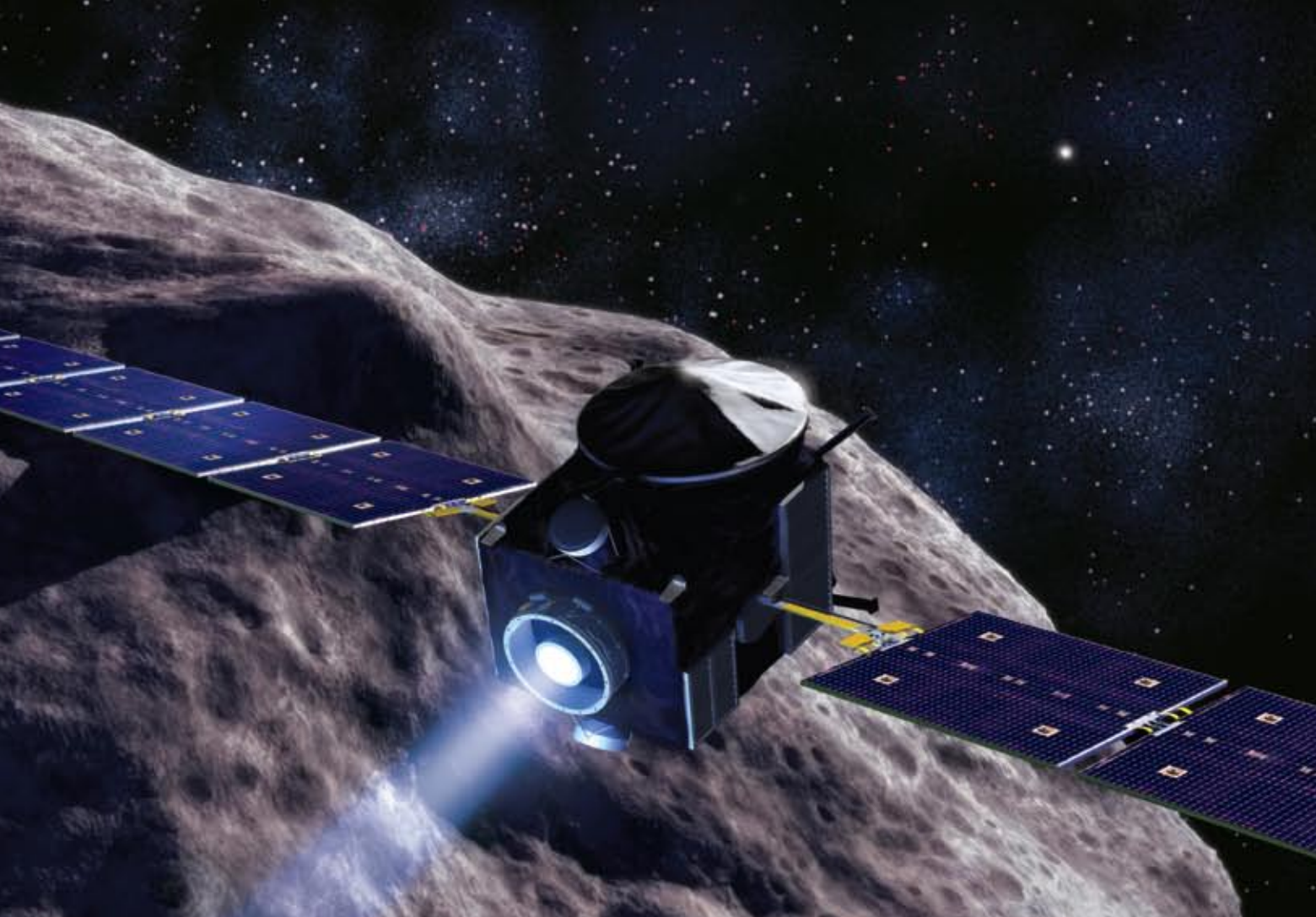
But the goals of the mission are not all that make this flight notable. Dawn, which took off in September 2007, is powered by a kind of space propulsion technology that is starting to take center stage for long-distance missions—a plasma rocket engine. Such engines, now being developed in several advanced forms, generate thrust by electrically producing and manipulating ionized gas propellants rather than by burn-

ing liquid or solid chemical fuels, as conventional rockets do.

Dawn's mission designers at the NASA Jet Propulsion Laboratory selected a plasma engine as the probe's rocket system because it is highly efficient, requiring only one tenth of the fuel that a chemical rocket motor would have needed to reach the asteroid belt. If project planners had chosen to install a traditional engine, the vehicle would have been able to reach either Vesta or Ceres, but not both.

Indeed, electric rockets, as the engines are also known, are quickly becoming the best option for sending probes to far-off targets. Recent successes made possible by electric propulsion include a visit by NASA's Deep Space 1 vehicle to a comet, a bonus journey that was made feasible by propellant that was left over after the spacecraft had ac-

PAT RAWLINGS SAIC



complished its primary goal. Plasma engines have also provided propulsion for an attempted landing on an asteroid by the Japanese Hayabusa probe, as well as a trip to the moon by the European Space Agency's SMART-1 spacecraft. In light of the technology's demonstrated advantages, deep-space mission planners in the U.S., Europe and Japan are opting to employ plasma drives for future missions that will explore the outer planets, search for extrasolar, Earth-like planets and use the void of space as a laboratory in which to study fundamental physics.

A Long Time Coming

Although plasma thrusters are only now making their way into long-range spacecraft, the technology has been under development for that purpose for some time and is already used for other tasks in space.

As early as the first decade of the 20th century, rocket pioneers speculated about using electricity to power spacecraft. But the late Ernst Stuhlinger—a member of Wernher von Braun's legendary team of German rocket scientists that spearheaded the U.S. space program—finally

turned the concept into a practical technology in the mid-1950s. A few years later engineers at the NASA Glenn Research Center (then known as Lewis) built the first operating electric rocket. That engine made a suborbital flight in 1964 on-board Space Electric Rocket Test 1, operating for half an hour before the craft fell back to Earth.

In the meantime, researchers in the former Soviet Union worked independently on concepts for electric rockets. Since the 1970s mission planners have selected the technology because it can save propellant while performing such tasks as maintaining the attitude and orbital position of telecommunications satellites in geosynchronous orbit.

Rocket Realities

The benefits afforded by plasma engines become most striking in light of the drawbacks of conventional rockets. When people imagine a ship streaking through the dark void toward a distant planet, they usually envision it trailing a long, fiery plume from its nozzles. Yet the truth is altogether different: expeditions to the outer solar system have been mostly rocketless affairs,

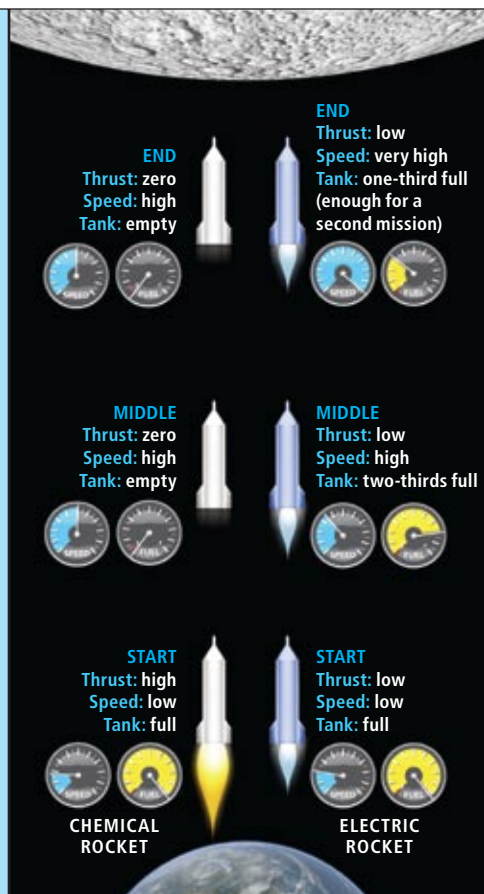
NASA'S DAWN SPACE PROBE, which is propelled by an electric rocket called an ion thruster, nears the asteroid Vesta in this artist's conception. Vesta is its initial survey target; the asteroid Ceres, its second destination, floats in the far distance in the image (bright spot at upper right). A conventional chemical rocket engine would be able to carry enough fuel to reach only one of these asteroids.

[COMPARISON]

Chemical vs. Electric Rockets

Chemical and electric propulsion systems are suited to different kinds of missions. Chemical rockets (*left*) produce large amounts of thrust quickly, so they can accelerate to high speeds rapidly, although they burn copious quantities of fuel to do so. These characteristics make them appropriate for relatively short-range trips.

Electric rockets (*right*), which use a plasma (ionized gas) as propellant, generate much less thrust, but their extremely frugal consumption of propellant allows them to operate for much longer periods. And in the frictionless environment of space, a small force applied over time can eventually achieve similarly high or greater speeds. These features make plasma rockets well equipped for deep-space missions to multiple targets.



motor would typically have no fuel left for braking. Such a probe would need the ability to fire its rocket so that it could slow enough to achieve orbit around its target and thus conduct extended scientific observations. Unable to brake, it would be limited to just a fleeting encounter with the object it aimed to study. Indeed, after a trip of more than nine years, New Horizons, a NASA deep-space probe launched in 2006, will get only a brief encounter of not more than a single Earth day with its ultimate object of study, the recently demoted “dwarf planet” Pluto.

The Rocket Equation

For those who wonder why engineers have been unable to come up with ways to send enough chemical fuel into space to avoid such difficulties for long missions, let me clarify the immense hurdles they face. The explanation derives from what is called the rocket equation, a formula used by mission planners to calculate the mass of propellant required for a given mission. Russian scientist Konstantin E. Tsiolkovsky, one of the fathers of rocketry and spaceflight, first introduced this basic formula in 1903.

In plain English, the rocket equation states the intuitive fact that the faster you throw propellant out from a spacecraft, the less you need to execute a rocket-borne maneuver. Think of a baseball pitcher (a rocket motor) with a bucket of baseballs (propellant) standing on a skateboard (a spacecraft). The faster the pitcher flings the balls rearward (that is, the higher the exhaust speed), the faster the vehicle will be traveling in the opposite direction when the last ball is thrown—or, equivalently, the fewer baseballs (less propellant) the pitcher would have to hurl to raise the skateboard’s speed by a desired amount at any given time. Scientists call this incremental increase of the skateboard’s velocity “delta-v.”

In more specific terms, the equation relates the mass of propellant required by a rocket to carry out a particular mission in outer space to two key velocities: the velocity at which the rocket’s exhaust will be ejected from the vehicle and the mission’s delta-v—how much the vehicle’s velocity will increase as a result of the exhaust’s ejection. Delta-v corresponds to the energy a craft must expend to alter its inertial motion and execute a desired space maneuver. For a given rocket technology (that is, one that produces a given rocket exhaust speed), the rocket equation translates the delta-v for a desired mission into the mass of propellant required to complete it.

[THE AUTHOR]

Edgar Y. Choueiri teaches astronautics and applied physics at Princeton University, where he also directs the Electric Propulsion and Plasma Dynamics Laboratory (<http://alfven.princeton.edu>) and the university’s Program in Engineering Physics. Aside from plasma propulsion research, he is working on mathematical techniques that could enable accurate recording and reproduction of music in three dimensions.



because most of the fuel is typically expended in the first few minutes of operation, leaving the spacecraft to coast the rest of the way to its goal. True, chemical rockets do launch all spacecraft from Earth’s surface and can make midcourse corrections. But they are impractical for powering deep-space explorations because they would require huge quantities of fuel—too much to be lifted into orbit practically and affordably. Placing a pound (0.45 kilogram) of anything into Earth orbit costs as much as \$10,000.

To achieve the necessary trajectories and high speeds for lengthy, high-precision journeys without additional fuel, many deep-space probes of the past have had to spend time—often years—detouring out of their way to planets or moons that provided gravitational kicks able to accelerate them in the desired direction (slingshot moves called gravity-assist maneuvers). Such circuitous flight paths limit missions to relatively small launch windows; only blasting off within a certain short time frame will ensure a precision swing past a cosmic body serving as a gravitational booster.

Even worse, after years of travel toward its destination, a vehicle with a chemical rocket

The delta-v metric can therefore be thought of as a kind of “price tag” of a mission, because the cost of conducting one is typically dominated by the cost of launching the needed propellant.

Conventional chemical rockets achieve only low exhaust velocities (three to four kilometers per second, or km/s). This feature alone makes them problematic to use. Also, the exponential nature of the rocket equation dictates that the fraction of the vehicle’s initial mass that is composed of fuel—the “propellant mass fraction”—grows exponentially with delta-v. Hence, the fuel needed for the high delta-v required for a deep-space mission could take up almost all the starting mass of the spacecraft, leaving little room for anything else.

Consider a couple of examples: To travel to Mars from low-Earth orbit requires a delta-v of about 4.5 km/s. The rocket equation says that a conventional chemical rocket would require that more than two thirds of the spacecraft’s mass be propellant to carry out such an interplanetary transfer. For more ambitious trips—such as expeditions to the outer planets, which have delta-v requirements that range from 35 to 70 km/s—chemical rockets would need to be more than 99.98 percent fuel. That configuration would leave no space for other hardware or useful payloads. As probes journey farther out into the solar system, chemical rockets become increasingly useless—unless engineers can find a way to significantly raise their exhaust speeds.

So far that goal has proved very difficult to accomplish because generating ultrahigh exhaust speeds demands extremely high fuel combustion temperatures. The ability to reach the needed temperatures is limited both by the amount of energy that can be released by known chemical reactions and by the melting point of the rocket’s walls.

The Plasma Solution

Plasma propulsion systems, in contrast, offer much greater exhaust speeds. Instead of burning chemical fuel to generate thrust, the plasma engine accelerates plasmas—clouds of electrically charged atoms or molecules—to very high velocities. A plasma is produced by adding energy to a gas, for instance, by radiating it with lasers, microwaves or radio-frequency waves or by subjecting it to strong electric fields. The extra energy liberates electrons from the atoms or molecules of the gas, leaving the latter with a positive charge and the former free to move freely in the gas, which makes the ionized gas a far

better electrical conductor than copper metal. Because plasmas contain charged particles, whose motion is strongly affected by electric and magnetic fields, application of electric or electromagnetic fields to a plasma can accelerate its constituents and send them out the back of a vehicle as thrust-producing exhaust. The necessary fields can be generated by electrodes and magnets, using induction by external antennas or wire coils, or by driving electric currents through the plasma.

The electric power for creating and accelerating the plasmas typically comes from solar panels that collect energy from the sun. But deep-space vehicles going past Mars must rely on nuclear power sources, because solar energy gets too weak at long distances from the sun. Today’s small robotic probes use thermoelectric devices heated by the decay of a nuclear isotope, but the more ambitious missions of the future would need nuclear fission (or even fusion) reactors. Any nuclear reactor would be activated only after the vessel reached a stable orbit at a safe distance from Earth. Its fuel would be secured in an inert state during liftoff.

Three kinds of plasma propulsion systems have matured enough to be employed on long-distance missions. The one in most use—and the kind powering Dawn—is the ion drive.

The Ion Drive

The ion engine, one of the more successful electric propulsion concepts, traces its roots to the ideas of American rocketry pioneer Robert H. Goddard, formed when he was still a graduate student at Worcester Polytechnic Institute a century ago. Ion engines are able to achieve exhaust velocities ranging from 20 to 50 km/s [see box on next page].

In its most common incarnation, the ion engine gets its electric power from photovoltaic panels. It is a squat cylinder, not much larger than a bucket, that is set astern. Inside the bucket, xenon gas from the propellant tank flows into an ionization chamber where an electromagnetic field tears electrons off the xenon gas atoms to create a plasma. The plasma’s positive ions are then extracted and accelerated to high speeds through the action of an electric field that is applied between two electrode grids. Each positive ion in the field feels the strong tug of the aft-mounted, negatively charged electrode and therefore accelerates rearward.

The positive ions in the exhaust leave a spacecraft with a net negative charge, which, if left to

EARLY HISTORY OF ELECTRIC ROCKETS

1903: Konstantin E. Tsiolkovsky derives the “rocket equation,” which is widely used to calculate fuel consumption for space missions. In 1911 he speculates that electric fields could accelerate charged particles to produce rocket thrust.

1906: Robert H. Goddard conceives of electrostatic acceleration of charged particles for rocket propulsion. He invents and patents a precursor to the ion engine in 1917.

1954: Ernst Stuhlinger figures out how to optimize the performance of the electric ion rocket engine.

1962: Work by researchers in the Soviet Union, Europe and the U.S. leads to the first published description of the Hall thruster, a more powerful class of plasma rocket.

1962: Adriano Ducati discovers the mechanism behind the magnetoplasmadynamic thruster, the most powerful type of plasma rocket.

1964: NASA’s SERT I spacecraft conducts the first successful flight test of an ion engine in space.

1972: The Soviet Meteor satellite carries out the initial spaceflight of a Hall thruster.

1999: NASA’s Jet Propulsion Laboratory’s Deep Space 1 demonstrates the first use of an ion engine as the main propulsion system on a spacecraft that escapes Earth’s gravitation from orbit.

ROBERT H. GODDARD, circa 1935

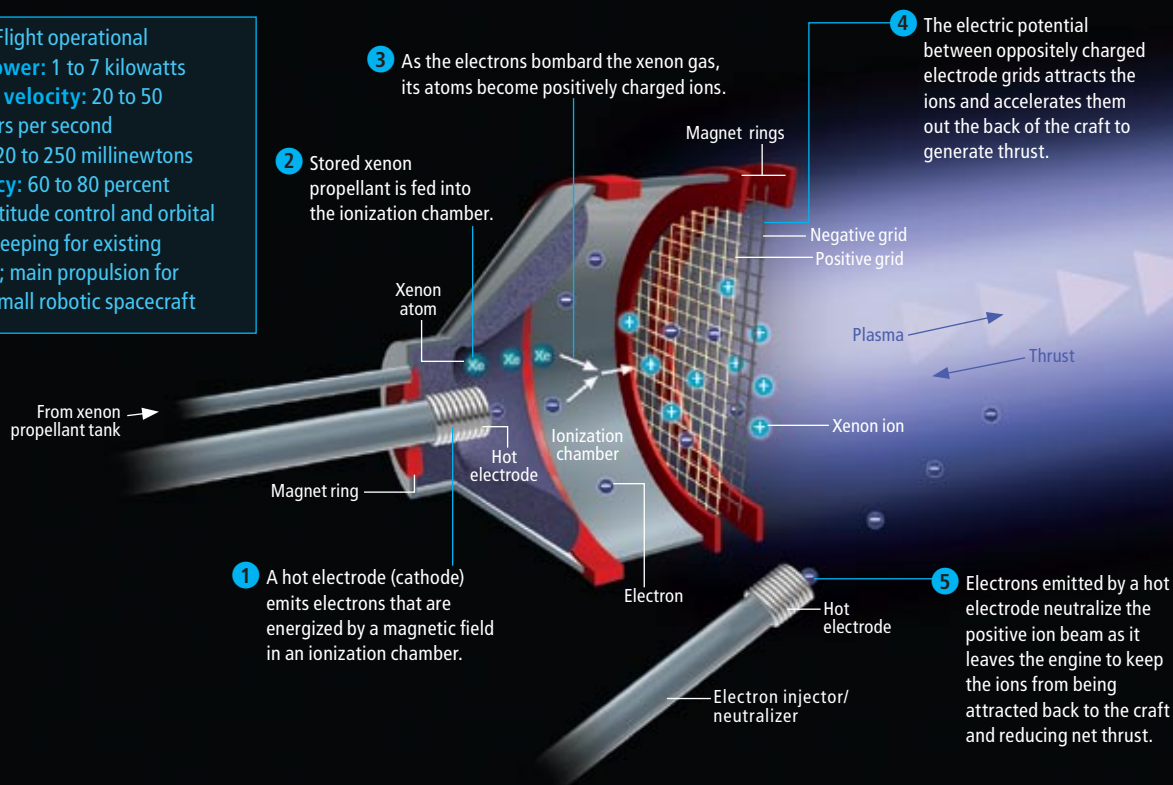


The Proven Plasma Propulsion Workhorse

This engine type creates a plasma propellant by bombarding a neutral gas with electrons emitted from a hot electric filament. The resulting positive ions are then extracted from the plasma and accelerated out

the back of the craft by an electric field that is created by applying a high voltage between two electrode grids. The ion exhaust generates thrust in the opposite direction.

Status: Flight operational
Input power: 1 to 7 kilowatts
Exhaust velocity: 20 to 50 kilometers per second
Thrust: 20 to 250 millinewtons
Efficiency: 60 to 80 percent
Uses: Attitude control and orbital station-keeping for existing satellites; main propulsion for current small robotic spacecraft



ION THRUSTER, which is 40 centimeters in diameter, was test-fired inside a laboratory vacuum chamber. Charged xenon atoms account for the blue color of the exhaust plume.

build up, would attract the ions back to the spacecraft, thus canceling out the thrust. To avoid this problem, an external electron source (a negative cathode or an electron gun) injects electrons into the positive flow to electrically neutralize it, which leaves the spacecraft neutral.

Dozens of ion drives are currently operating on commercial spacecraft—mostly communications satellites in geosynchronous orbit for orbital “station-keeping” and attitude control. They were selected because they save millions of dollars per spacecraft by greatly shrinking the mass of propellant that would be required for chemical propulsion.

At the end of the 20th century, Deep Space 1 became the world’s first spacecraft using an electric propulsion system to escape Earth’s gravitation from orbit. The probe then accelerated by about 4.3 km/s, while consuming less than 74 kilograms of xenon propellant (about the mass of an untapped beer keg), to fly through the dusty tail of the comet Borrelly. This is the highest velocity increment gained via propulsion (as opposed to gravity assists) by any spacecraft to

date. Dawn should soon break that record by adding 10 km/s to its velocity. Engineers at the Jet Propulsion Laboratory have recently demonstrated ion drives able to function flawlessly for more than three years of continuous operation.

A plasma rocket’s performance is determined not only by the speed of the exhaust particles but also by its thrust density, which is the amount of thrust force an engine produces per unit area of its exhaust aperture. Ion engines and similar electrostatic thrusters suffer from a major shortcoming, called space-charge limitation, that severely reduces their thrust density: as the positive ions pass between the electrostatic grids in an ion engine, a positive charge inevitably builds up in this region. This buildup limits the attainable electric field to drive the acceleration.

Because of this phenomenon, Deep Space 1’s ion engine produces a thrust force that is roughly equivalent to the weight of a single sheet of paper—hardly the thundering rocket engine of sci-fi movies and more akin to a car that takes two days to accelerate from zero to 60 miles per

hour. As long as one is willing to wait long enough (typically, many months), though, these engines can eventually attain the high delta-vs needed for distant journeys. That feat is possible because in the vacuum of space, which offers no resistance, even a tiny push, if constantly applied, will lead to high propulsion speeds.

The Hall Thruster

A plasma propulsion system called the Hall thruster [see box at right] avoids the space-charge limitation and can therefore accelerate a vessel to high speeds more quickly (by way of its greater thrust density) than a comparably sized ion engine can. This technology has been gaining acceptance in the West since the early 1990s, after three decades of steady development in the former Soviet Union. The Hall thruster will soon be ready to take on long-range missions.

The system relies on a fundamental effect discovered in 1879 by Edwin H. Hall, then a physics graduate student at Johns Hopkins University. Hall showed that when electric and magnetic fields are set perpendicular to each other inside a conductor, an electric current (called the Hall current) flows in a direction that is perpendicular to both fields.

In a Hall thruster a plasma is produced when an electric discharge between an internal positive anode and a negative cathode situated outside the device tears through a neutral gas inside the device. The resulting plasma fluid is then accelerated out of the cylindrical engine by the Lorentz force, which results from the interaction of an applied radial magnetic field and an electric current (in this case, the Hall current) that flows in an azimuthal direction—that is, in a circular “orbit” around the central anode. The Hall current is caused by the electron’s motion in the magnetic and electric fields. Depending on the available power, exhaust velocities can range from 10 to more than 50 km/s.

This form of electric rocket avoids a space-charge buildup by accelerating the entire plasma (of both positive ions and negative electrons), with the result that its thrust density and thus its thrust force (and so its potential delta-v) is many times that of an ion engine of the same size. More than 200 Hall thrusters have been flown on satellites in Earth orbit. And it was a Hall thruster that the European Space Agency used to efficiently propel its SMART-1 spacecraft frugally to the moon.

Engineers are now trying to scale up today’s

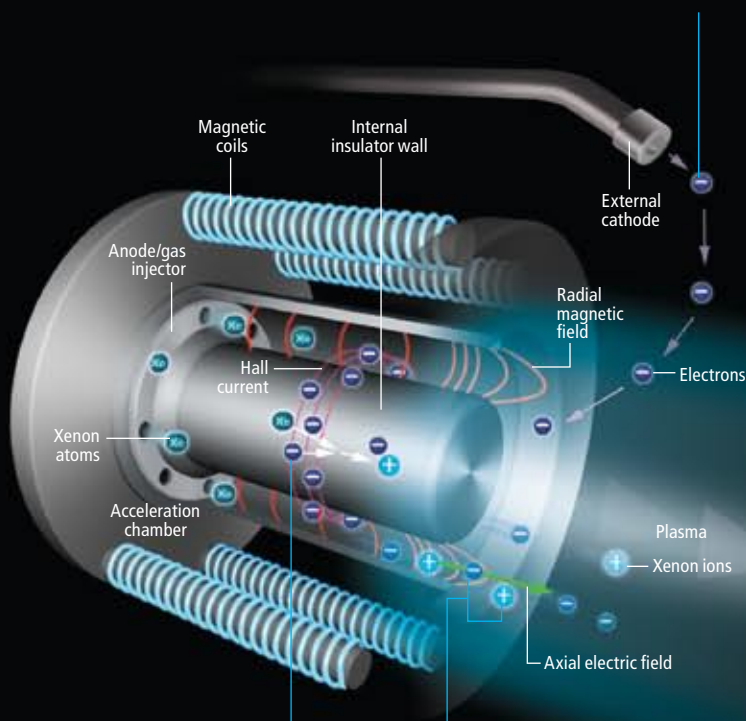
[HALL THRUSTER]

The Latest Plasma Engine Contender

This device generates propulsion by crossing a so-called Hall current and a radial magnetic field, which causes electrons to circle around the device’s axis. These electrons tear electrons from xenon atoms, producing xenon ions, and an electric field parallel to the axis accelerates the ions downstream. The density of propulsive force produced by a Hall thruster is greater than that of an ion engine because its exhaust contains both positive ions and electrons, which avoids the buildup of positive charge that can limit the strength of an accelerating electric field.

1 An electric potential established between an external negative cathode and internal positive anode creates a mostly axial electric field inside the acceleration chamber.

2 As the cathode heats up, it emits electrons. Some of the electrons drift upstream toward the anode. When the electrons enter the chamber, a radial magnetic field and the axial electric field cause them to whirl around the axis of the thruster as a “Hall current.”



3 Xenon gas propellant feeds through the positive anode injector into the annular acceleration chamber, where the whirling electrons collide with the xenon atoms, turning them into positive ions.

4 The plasma (containing both positive ions and electrons) is accelerated sternward by the electromagnetic forces resulting from the interaction between the predominantly radial magnetic field and the Hall current.

Status: Flight operational

Input power: 1.35 to 10 kilowatts

Exhaust velocity: 10 to 50 kilometers per second

Thrust: 40 to 600 millinewtons

Efficiency: 45 to 60 percent

Uses: Satellite attitude control and station-keeping; used as main propulsion for medium-size robotic spacecraft

\$10,000

is roughly what it costs to send a pound (0.45 kilogram) of payload into Earth orbit with conventional rocket boosters. This high price tag is one reason engineers go to great lengths to shave as much mass from spacecraft as is feasible. The fuel and its storage tank are the heaviest parts of a vehicle powered by a chemical rocket.

rather small Hall thrusters so that they can handle higher amounts of power to generate greater exhaust speeds and thrust levels. The work also aims to extend their operating lifetimes to the multiyear durations needed for deep-space exploration.

Scientists at the Princeton Plasma Physics Laboratory have taken a step toward these goals by implanting segmented electrodes in the walls of a Hall thruster. The electrodes shape the internal electric field in a way that helps to focus the plasma into a thin exhaust beam. This design reduces the useless nonaxial component of thrust and improves the system's operating lifetime by keep-

ing the plasma beam away from the thruster walls. German engineers have achieved similar results using specially shaped magnetic fields. Researchers at Stanford University have meanwhile shown that lining the walls with tough, synthetic-polycrystalline diamond substantially boosts the device's resistance to plasma erosion. Such improvements will eventually make Hall thrusters suitable for deep-space missions.

Next-Generation Thruster

One way to further raise the thrust density of plasma propulsion is to increase the total amount of plasma that is accelerated in the engine. But as the plasma density in a Hall thruster is raised, electrons collide more frequently with atoms and ions, which makes it more difficult for the electrons to carry the Hall current needed for acceleration. An alternative known as the magnetoplasmadynamic thruster (MPDT) allows for a denser plasma by forgoing the Hall current in favor of a current component that is mostly aligned with the electric field [see box at left] and far less prone than the Hall current to disruption by atomic collisions.

In general, an MPDT consists of a central cathode sitting within a larger cylindrical anode. A gas, typically lithium, is pumped into the annular space between the cathode and the anode. There it is ionized by an electric current flowing radially from the cathode to the anode. This current induces an azimuthal magnetic field (one that encircles the central cathode), which interacts with the same current that induced it to generate the thrust-producing Lorentz force.

A single MPD engine about the size of an average household pail can process about a million watts of electric power from a solar or nuclear source into thrust (enough to energize more than 10,000 standard lightbulbs), which is substantially larger than the maximum power limits of ion or Hall thrusters of the same size. An MPDT can produce exhaust velocities from 15 to 60 km/s. It truly is the little engine that could.

This design also offers the advantage of throttling; its exhaust speed and thrust can be easily adjusted by varying the electric current level or the flow rate of the propellant. Throttling allows a mission planner to alter a spacecraft's engine thrust and exhaust velocity as needed to optimize its trajectory.

Intensive research on mechanisms that hamper the performance and lifetimes of MPD devices, such as electrode erosion, plasma instabilities

[MAGNETOPLASMADYNAMIC THRUSTER]

The Future of Plasma Propulsion

An MPDT relies on the Lorentz electromagnetic force to accelerate the plasma to produce thrust. The Lorentz force (*green arrows*), which is mainly along the axis, is created by the interaction of a mostly radial electric current pattern (*red lines*) with a concentric magnetic field (*blue circle*).

Status: Flight-tested but not yet operational

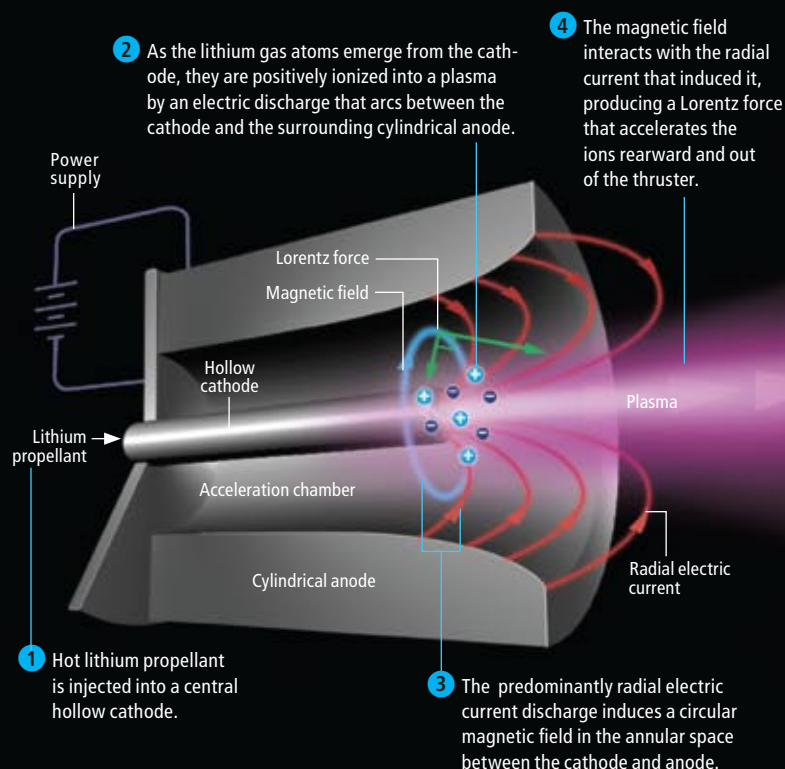
Input power: 100 to 500 kilowatts

Exhaust velocity: 15 to 60 kilometers per second

Thrust: 2.5 to 25 newtons

Efficiency: 40 to 60 percent

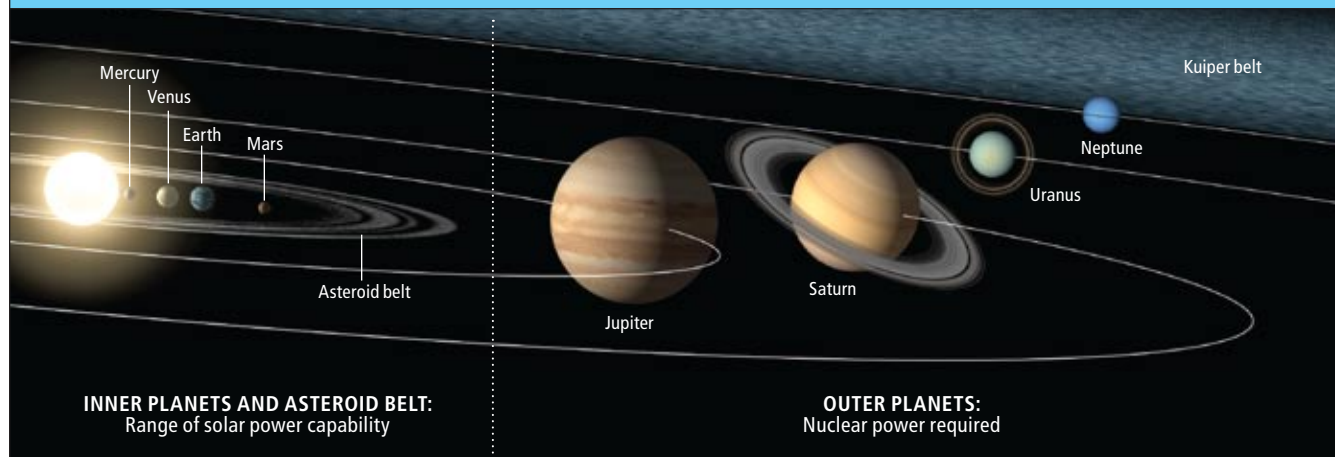
Uses: Main propulsion for heavy cargo and piloted spacecraft; under development



Solar and Nuclear Energy for Electric Rockets

For trips to the inner solar system, where the sun's rays are strong, sufficient electric power can be provided to plasma rocket engines by solar cells. But trips to the outer planets of the solar system would

generally require nuclear power sources. A large, heavy craft would need a nuclear reactor, but a smaller, lighter one might get by with a thermoelectric power-generation device heated by the decay of radioisotopes.



and power dissipation in the plasma, has led to new, high-performance engines that rely on lithium and barium vapors for propellants. These elements ionize easily, yield lower internal energy losses in the plasma and help to keep the cathode cooler. The adoption of these liquid-metal propellants and an unusual cathode design that contains channels that alter how the electric current interacts with its surface has resulted in substantially less erosion of the cathode. These innovations are leading to more reliable MPDTs.

A team of academic and NASA researchers has recently completed the design of a state-of-the-art lithium-fed MPDT called α^2 , which could potentially drive a nuclear-powered vessel hauling heavy cargo and people to the moon and Mars as well as robotic missions to the outer planets.

The Tortoise Wins

Ion, Hall and MPD thrusters are but three variants of electric plasma rocket technology, albeit the most mature. During the past few decades researchers have developed many other promising related concepts to various degrees of readiness. Some are pulsed engines that operate intermittently; others run continuously. Some generate plasmas through electrode-based electric discharge; others use coil-based magnetic induction or antenna-generated radiation. The mechanisms they apply to accelerate plasmas vary as well: some use Lorentz forces; others accelerate the plasmas by entraining them in magnetically produced current sheets or in traveling electromagnetic waves. One type even aims to exhaust

the plasma through invisible “rocket nozzles” composed of magnetic fields.

In all cases, plasma rockets will get up to speed more slowly than conventional rockets. And yet, in what has been called the “slower but faster paradox,” they can often make their way to distant destinations more quickly by ultimately reaching higher spacecraft velocities than standard propulsion systems can using the same mass of propellant. They thus avoid time-consuming detours for gravity boosts. Much as the fabled slow and steady tortoise beats out the intermittently sprinting hare, in the marathon flights that will become increasingly common in the coming era of deep-space exploration, the tortoise wins.

So far the most advanced designs could impart a delta-v of 100 km/s—much too slow to take a spacecraft to the far-off stars but plenty enough to visit the outer planets in a reasonable amount of time. One particularly exciting deep-space mission that has been proposed would return samples from Saturn’s largest moon, Titan, which space scientists believe has an atmosphere that is very similar to Earth’s eons ago.

A sample from Titan’s surface would offer researchers a rare chance to search for signs of chemical precursors to life. The mission would be impossible with chemical propulsion. And, with no in-course propulsion, the journey would require multiple planetary gravity assists, adding more than three years to the total trip time. A probe fitted with “the little plasma engine that could” would be able to do the job in a significantly shorter period.

MORE TO EXPLORE

Benefits of Nuclear Electric Propulsion for Outer Planet Exploration. G. Woodcock et al. American Institute of Aeronautics and Astronautics, 2002.

Electric Propulsion. Robert G. Jahn and Edgar Y. Choueiri in *Encyclopedia of Physical Science and Technology*. Third edition. Academic Press, 2002.

A Critical History of Electric Propulsion: The First 50 Years (1906–1956). Edgar Y. Choueiri in *Journal of Propulsion and Power*, Vol. 20, No. 2, pages 193–203; 2004.

Physics of Electric Propulsion. Robert G. Jahn. Dover Publications, 2006.

Fundamentals of Electric Propulsion: Ion and Hall Thrusters. Dan M. Goebel and Ira Katz. Wiley, 2008.

Sculpting the BRAIN

New studies are revealing how the brain's convolutions take shape—findings that could aid the diagnosis and treatment of autism, schizophrenia and other mental disorders

BY CLAUS C. HILGETAG AND HELEN BARBAS

KEY CONCEPTS

- The cerebral cortex is the structure that gives the organ its convoluted surface. It is involved with high-level processing of our perceptions, thoughts, emotions and actions.
- Intricate folding permits the expansive cortex to fit inside a skull with limited surface area.
- Recent discoveries have shown that mechanical tension between neurons creates the hills and valleys of the cortex.
- The cortical landscape differs between healthy people and individuals with brain disorders that originate during development, such as autism. These shape differences suggest that connections between brain regions of affected individuals also depart from the norm.
—The Editors

One of the first things people notice about the human brain is its intricate landscape of hills and valleys. These convolutions derive from the cerebral cortex, a two- to four-millimeter-thick mantle of gelatinous tissue packed with neurons—sometimes called gray matter—that mediates our perceptions, thoughts, emotions and actions. Other large-brained mammals such as whales, dogs and our great ape cousins have a corrugated cortex, too—each with its own characteristic pattern of convolutions. But small-brained mammals and other vertebrates have relatively smooth brains. The cortex of large-brained mammals expanded considerably over the course of evolution—much more so than the skull. Indeed, the surface area of a flattened human cortex—equivalent to that of an extra-large pizza—is three times larger than the inner surface of the braincase. Thus, the only way the cortex of humans and other brainy species can fit into the skull is by folding.

This folding is not random, as in a crumpled piece of paper. Rather it exhibits a pattern that is consistent from person to person. How does this folding occur in the first place? And what, if anything, can the resulting topography reveal about brain function? New research indicates that a network of nerve fibers physically pulls the pliable cortex into shape during development and holds it in place throughout life. Disturbances to

this network during development or later, as a result of a stroke or injury, can have far-reaching consequences for brain shape and neural communication. These discoveries could therefore lead to new strategies for diagnosing and treating patients with certain mental disorders.

Internal Forces

Scientists have pondered the brain's intricate form for centuries. In the early 1800s German physician Franz Joseph Gall proposed that the shape of a person's brain and skull spoke volumes about that individual's intelligence and personality—a theory known as phrenology. This influential, albeit scientifically unsupported, idea led to the collection and study of “criminal,” “degenerate” and “genius” brains. Then, in the latter part of the 19th century, Swiss anatomist Wilhelm His posited that the brain develops as a sequence of events guided by physical forces. British polymath D'Arcy Thompson built on that foundation, showing that the shapes of many structures, biological and inanimate, result from physical self-organization.

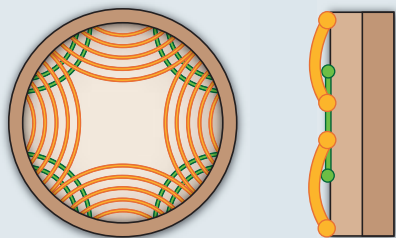
Provocative though they were, these early suppositions eventually faded from view. Phrenology became known as a pseudoscience, and mod-

TORTUOUS FOLDING enables the human brain's oversize cerebral cortex to fit into the skull.

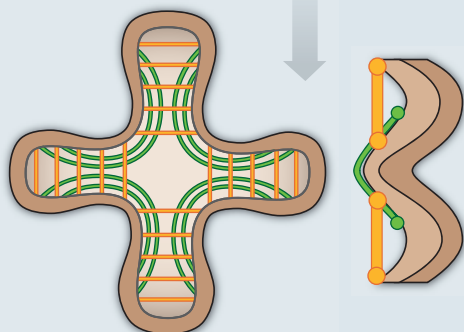


FOLDING THE BRAIN

The wrinkles of the brain's outermost layer, the cerebral cortex, arise in the womb. Studies suggest that mechanical forces generated by neurons connecting different brain regions drive the folding, as depicted in the simplified representations of the cortex (below left).

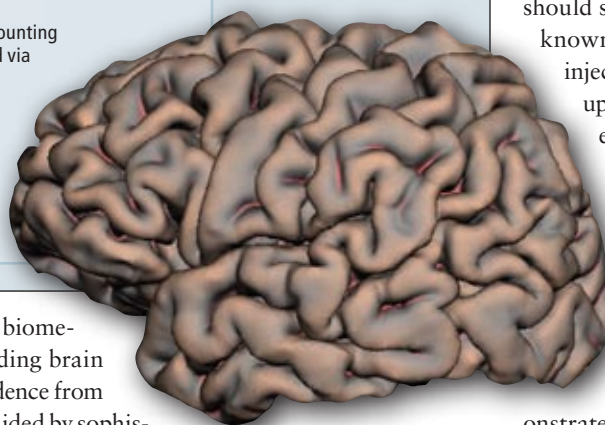
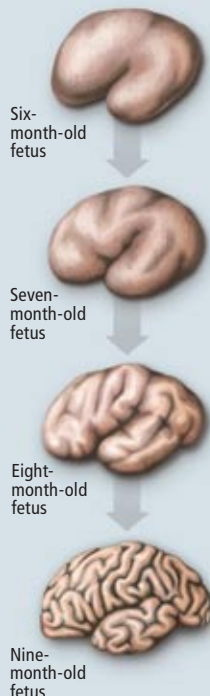


During the first 25 weeks of fetal development, the cortex remains relatively smooth while the emerging neurons send out fibers (colored lines) to connect with neurons in other regions of the brain, where they become tethered.



As the cortex continues to grow, mounting tension between regions connected via numerous fibers (orange) begins to draw them together, producing a bulge, or gyrus, between them. Weakly connected regions (green) drift apart, creating a valley, or sulcus.

The folding is mostly complete by the time of birth.



ern genetic theories eclipsed the biomechanical approach to understanding brain structure. Recently, however, evidence from novel brain-imaging techniques, aided by sophisticated computational analyses, has lent fresh support to some of those 19th-century notions.

Hints that His and Thompson were on the right track with their ideas about physical forces shaping biological structures surfaced in 1997. Neurobiologist David Van Essen of Washington University in St. Louis published a hypothesis in *Nature* in which he suggested that the nerve fibers that link different regions of the cortex, thereby enabling them to communicate with one another, produce small tension forces that pull at this gelatinous tissue. In a human fetus, the cortex starts out smooth and mostly remains

that way for the first six months of development. During that time, the newborn neurons send out their spindly fibers, or axons, to hook up with the receptive components, or dendrites, of target neurons in other regions of the cortex. The axons then become tethered to the dendrites. As the cortex expands, the axons grow ever more taut, stretching like rubber bands. Late in the second trimester, while neurons are still emerging, migrating and connecting, the cortex begins to fold. By the time of birth the cortex has more or less completed development and attained its characteristically wrinkled form.

Van Essen argued that two regions that are strongly linked—that is, connected via numerous axons—are drawn closer together during development by way of the mechanical tension along the tethered axons, producing an outward bulge, or gyrus, between them. A pair of weakly connected regions, in contrast, drifts apart, becoming separated by a valley, or sulcus.

Modern techniques for tracing neural pathways have made it possible to test the hypothesis that the communication system of the cortex is also responsible for shaping the brain. According to a simple mechanical model, if each axon pulls with a small amount of force, the combined force of axons linking strongly connected areas should straighten the axon paths. Using a tool

known as retrograde tracing, in which a dye injected in a small area of the cortex is taken up by the endings of axons and transported backward to the parent cell body, it is possible to show which regions send axons to the injection site. Furthermore, the method can reveal both how dense the connections of an area are and what shapes their axon paths take. Our retrograde tracing studies of a large number of neural connections in the rhesus macaque have demonstrated that, as predicted, most connections follow straight or slightly curved paths. Moreover, the denser the connections are, the straighter they tend to run.

The sculpting power of neural connections is particularly evident in the shape differences between language regions in the left and right hemispheres of the human brain [see “Specializations of the Human Brain,” by Norman Geschwind; *SCIENTIFIC AMERICAN*, September 1979]. Take, for instance, the form of the Sylvian fissure—a prominent sulcus that separates the frontal and posterior language regions. The fissure on the left side of the brain is quite a bit shall-

lower than the one on the right. The asymmetry seems to be related to the anatomy of a large fiber bundle called the arcuate fascicle, which travels around the fissure to connect the frontal and posterior language regions. Based on this observation and the fact that the left hemisphere is predominantly responsible for language in most people, we postulated in a paper published in 2006 that the arcuate fascicle on the left is denser than the one on the right. A number of imaging studies of the human brain have confirmed this asymmetrical fiber density. In theory, then, the larger fiber bundle should have a greater pulling strength and therefore be straighter than the bundle on the right side. This hypothesis has yet to be tested, however.

From Macro to Micro

Mechanical forces shape more than just the large-scale features of the cerebral cortex. They also have an effect on its layered structure. The cortex is made of horizontal tiers of cells, stacked up as in a multilayered cake. Most areas have six layers, and individual layers in those areas vary in thickness and composition. For example, the regions of the cortex that govern the primary senses have a thick layer 4, and the region that controls voluntary motor functions has a thick layer 5. Meanwhile the association areas of the cortex—which underlie thinking and memory, among other things—have a thick layer 3.

Such variations in laminar structure have been used to divide the cortex into specialized areas for more than 100 years, most famously by German anatomist Korbinian Brodmann, who created a map of the cortex that is still in use today. Folding changes the relative thickness of the layers, as would happen if one were to bend a stack of sponges. In the gyri, the top layers of the cortex are stretched and thinner, whereas in the sulci, the top layers are compressed and thicker. The relations are reversed in the deep layers of the cortex.

Based on these observations, some scientists had suggested that whereas the shapes of layers and neurons change as they are stretched or compressed, the total area of the cortex and the number of neurons it comprises are the same. If so, thick regions (such as the deep layers of gyri) should contain fewer neurons than thin regions of the cortex. This isometric model, as it is known, assumes that during development neurons first migrate to the cortex and then the cortex folds. For an analogy, imagine folding a bag of rice. The shape of the bag chang-

THE AUTHORS



Claus C. Hilgetag is associate professor of neuroscience at Jacobs University Bremen, a new German research university that he joined as part of the founding faculty in 2001. His research in computational neuroscience centers on brain connectivity. **Helen Barbas** is a professor at Boston University, where her research focuses on the prefrontal cortex. She received her Ph.D. from McGill University. In addition to her deep interest in patterns in brain circuits, Barbas, an avid gardener, is fascinated by patterns in nature.

PHRENOLOGY REVISITED?

Popular in the 19th century, phrenology was the practice of reading the form of the skull for insights into an individual's personality traits and mental ability. Practitioners believed that the skull's bumps and depressions arose from the shape of the brain, with each area corresponding to a different mental faculty. The discipline was eventually dismissed as pseudoscience. But neuroscientists have since determined that brain shape (though not skull shape) may generally correlate with mental function and dysfunction. They have yet to identify specific patterns of difference between the brains of normal individuals and those who are geniuses or criminals, however.



es, but its capacity and the number of grains are the same before and after folding.

Our investigations into the density of neurons in areas of the prefrontal cortex in rhesus macaques reveal that the isometric model is wrong, though. Using estimates based on representative samples of frontal cortex, we determined that the deep layers of gyri are just as densely populated with neurons as the deep layers of sulci. And because the deep layers of gyri are thicker, there are actually more neurons under a unit area in gyri than in sulci.

Our discovery hinted that the physical forces that mold the gyri and sulci also influence neuronal migration. Developmental studies in humans have bolstered this suggestion. Rather than occurring sequentially, the migration of neurons to the cortex and the folding of the cortex partially overlap in time. Consequently, as the cortex folds, the resulting stretching and compressing of layers may well affect the passage of newly born neurons that migrate into the cortex late in development, which in turn would affect the composition of the cortex.

Moreover, the shapes of individual neurons differ depending on where in the cortex they reside. Neurons situated in the deep layers of gyri, for example, are squeezed from the sides and appear elongated. In contrast, neurons located in the deep layers of sulci are stretched and look flattened. The shapes of these cells are consistent with having been modified by mechanical forces as the cortex folded. It will be an intriguing challenge to figure out whether such systematic differences in the shapes of neurons in gyri and sulci also affect their function.

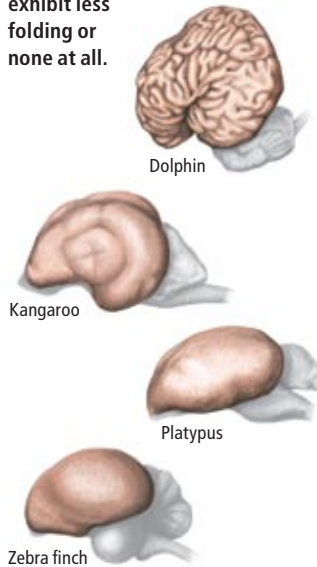
Our computer simulations suggest that they do: for example, because the cortical sheet is much thicker in the gyri than in the sulci, signals impinging on the dendrites of neurons at the bottom of a gyrus must travel a longer distance to the cell body than do signals impinging on the dendrites of neurons at the bottom of a sulci. Researchers can test the effect of these physical differences on the function of neurons by recording the activity of individual neurons across the undulating cortical landscape—work that, to our knowledge, has yet to be undertaken.

An Ill Influence?

To fully understand the relation between form and function, scientists will need to look at an abundance of brains. The good news is that now that we can observe the living human brain using noninvasive imaging techniques, such as struc-

OTHER BRAINSCAPES

Humans and other large mammals have elaborately folded cerebral cortices. But other vertebrates exhibit less folding or none at all.



Brains not shown to scale

tural magnetic resonance imaging, and reconstruct it in three dimensions on computers, we are able to collect images of a great many brains—far more than are featured in any of the classical collections of brains obtained after death. Researchers are systematically studying these extensive databases using sophisticated computer programs to analyze brain shape. One of the key findings from this research is that clear differences exist between the cortical folds of healthy people and those of patients with mental diseases that start in development, when neurons, connections and convolutions form. The mechanical relation between fiber connections and convolutions may explain these deviations from the norm.

Research into this potential link is still in its early phases. But over the past couple of years, several research teams have reported that the brains of schizophrenic patients exhibit reduced cortical folding overall relative to the brains of normal people. The findings are controversial, because the location and type of the folding aberrations vary considerably from person to person. We can say with certainty, however, that brain shape generally differs between schizophrenics and healthy people. Experts have often

attributed schizophrenia to a perturbed neurochemical balance. The new work suggests that there is additionally a flaw in the circuitry of the brain's communication system, although the nature of the flaw remains unknown.

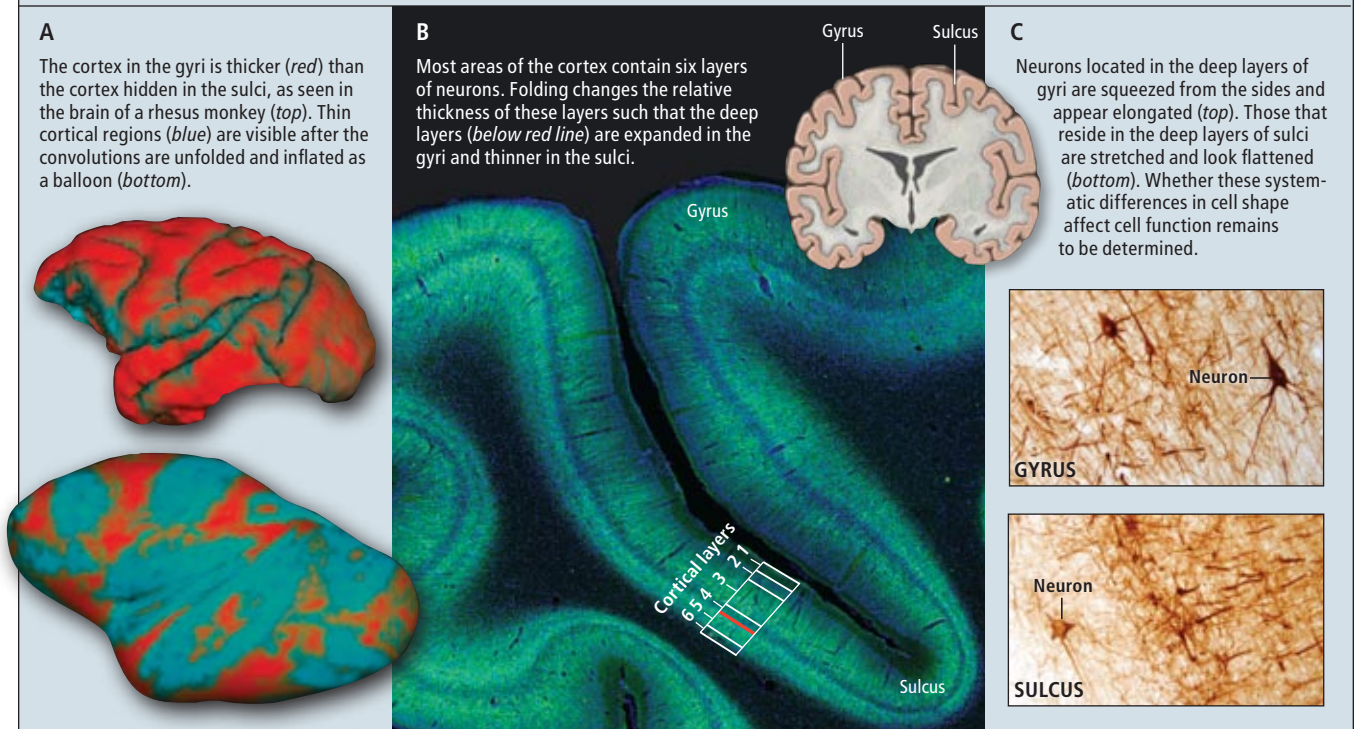
People diagnosed with autism also exhibit abnormal cortical convolutions. Specifically, some of their sulci appear to be deeper and slightly out of place as compared with those of healthy subjects. In light of this finding, researchers have begun to conceive of autism as a condition that arises from the miswiring of the brain. Studies of brain function support that notion, showing that in autistic people, communication between nearby cortical areas increases, whereas communication between distant areas decreases. As a result, these patients have difficulties ignoring irrelevant things and shifting their attention when it is appropriate to do so.

Mental disorders and learning disabilities can also be associated with aberrations in the composition of the cortical layers. For instance, in the late 1970s neurologist Albert Galaburda of Harvard Medical School found that in dyslexia the pyramidal neurons that form the chief communication system of the cerebral cortex

[EFFECTS OF MECHANICAL FORCES]

WHAT LIES BENEATH

Physical forces mold multiple aspects of the cerebral cortex, from large-scale features such as the thickness of the gyri and sulci (a) to the structure of the layers within the cortex (b) and the shapes of the neurons themselves (c).



JEN CHRISTIANSEN (animal brains and brain cross section), COURTESY OF HELEN BARBAS (micrographs)

are shifted from their normal position in the layers of the language and auditory areas of the frontal cortex. Schizophrenia, too, may leave imprints on the cortical architecture: some frontal areas of the cortex in affected individuals are aberrant in their neural density. Abnormal distribution of neurons in cortical layers disrupts their pattern of connections, which ultimately impairs the fundamental function of the nervous system in communication. Researchers are just beginning to probe structural abnormalities of the cortex in people with autism, which may further elucidate this puzzling condition.

More studies are needed to ascertain whether other neurological diseases that originate during development also bring about changes in the number and positions of neurons in the cortical layers. Thinking of schizophrenia and autism as disorders that affect neural networks, as opposed to localized parts of the brain, might lead to novel strategies for diagnosis and treatment. For instance, patients who have these conditions might profit from performing tasks that engage different parts of the brain, just as dyslexics benefit from using visual and multimodal aids in learning.

Modern neuroimaging methods have also enabled scientists to test the phrenological notion that cortical convolutions or the amount of gray matter in different brain regions can reveal a person's talents. Here, too, linking form and function is fraught with difficulty. The connection is clearest in people who routinely engage in well-defined coordinated mental and physical exercise.

Professional musicians offer one such example. These individuals, who need to practice extensively, systematically differ from nonmusicians in motor regions of the cortex that are involved in the control of their particular instruments. Still, clear folding patterns that distinguish broader mental talents remain elusive.

Vexing Variation

We still have much to unravel. For one, we do not yet understand how individual gyri attain their specific size and shape, any more than we understand the developmental underpinnings of variation in ear or nose shape among individuals. Variation is a very complex problem. Computational models that simulate the diversity of physical interactions between neurons during cortical development may throw light on this question in the future. So far, however, the models are very preliminary because of the complex-

ity of the physical interactions and the limited amount of developmental data available.

Scholars are also keen to know more about how the cortex develops. Topping our wish list is a detailed timetable of the formation of the many different connections that make up its extensive communication system. By marking neurons in animals, we will be able to determine when different parts of the cortex develop in the womb, which, in turn, will enable scientists to experimentally modify the development of distinct layers or neurons. Information about the sequence of development will help reveal events that result in abnormal brain morphology and function. The range of neurological diseases with vastly different symptoms—such as those seen in schizophrenia, autism, Williams syndrome, childhood epilepsy, and other disorders—may be the result of pathology arising at different times in development and variously affecting regions, layers and sets of neurons that happen to be emerging, migrating or connecting when the process goes awry.

To be sure, mechanical forces are not alone in modeling the brain. Comparisons of brain shape have demonstrated that brains of closely related people are more similar to one another than are brains of unrelated people, indicating that genetic programs are at work, too. Perhaps genetic processes control the timing of development of the cortex, and simple physical forces shape the brain as nerve cells are born, migrate and interconnect in a self-organizing fashion. Such a combination may help explain the remarkable regularity of the major convolutions among individuals, as well as the diversity of small convolutions, which differ even in identical twins.

Many of the current concepts about the shape of the brain have come full circle from ideas first proposed more than a century ago, including the notion of a link between brain shape and brain function. Systematic comparisons of brain shape across populations of normal subjects and patients with brain disorders affirm that the landscape of the brain does correlate with mental function and dysfunction.

But even with the advanced imaging methods for measuring brains, experts still cannot recognize the cortex of a genius or a criminal when they see one. New models of cortex folding that combine genetics and physical principles will help us integrate what we know about morphology, development and connectivity so that we may eventually unlock these and other secrets of the brain. ■



FORM AND FUNCTION: People with autism and other mental disorders that arise during fetal development have cortical folds that differ from those of healthy individuals. The composition of their cortical layers may also exhibit aberrations.

➔ MORE TO EXPLORE

On Growth and Form. D'Arcy Wentworth Thompson. Reprinted edition. Cambridge University Press, 1961.

A Tension-Based Theory of Morphogenesis and Compact Wiring in the Central Nervous System. David C. Van Essen in *Nature*, Vol. 385, pages 313–318; January 23, 1997.

Postcards from the Brain Museum: The Improbable Search for Meaning in the Matter of Famous Minds. Brian Burrell. Broadway, 2005.

Role of Mechanical Factors in the Morphology of the Primate Cerebral Cortex. Claus C. Hilgetag and Helen Barbas in *PLoS Computational Biology*, Vol. 2, No. 3, page e22; March 24, 2006. Available free at www.ploscompbiol.org

THE GREENHOUSE HAMBURGER

Producing beef for the table has a surprising environmental cost: it releases prodigious amounts of heat-trapping greenhouse gases

By Nathan Fiala

KEY CONCEPTS

- Pound for pound, beef production generates greenhouse gases that contribute more than 13 times as much to global warming as do the gases emitted from producing chicken. For potatoes, the multiplier is 57.
- Beef consumption is rising rapidly, both as population increases and as people eat more meat.
- Producing the annual beef diet of the average American emits as much greenhouse gas as a car driven more than 1,800 miles.

—The Editors

Most of us are aware that our cars, our coal-generated electric power and even our cement factories adversely affect the environment. Until recently, however, the foods we eat had gotten a pass in the discussion. Yet according to a 2006 report by the United Nations Food and Agriculture Organization (FAO), our diets and, specifically, the meat in them cause more greenhouse gases—carbon dioxide (CO₂), methane, nitrous oxide, and the like—to spew into the atmosphere than either transportation or industry. (Greenhouse gases trap solar energy, thereby warming the earth's surface. Because gases vary in greenhouse potency, every greenhouse gas is usually expressed as an amount of CO₂ with the same global-warming potential.)

The FAO report found that current production levels of meat contribute between 14 and 22 percent of the 36 billion tons of “CO₂-equivalent” greenhouse gases the world produces every year. It turns out that producing half a pound of hamburger for someone's lunch—a patty of meat the size of two decks of cards—releases as much greenhouse gas into the atmosphere as driving a 3,000-pound car nearly 10 miles.

In truth, every food we consume, vegetables and fruits included, incurs hidden environmen-

tal costs: transportation, refrigeration and fuel for farming, as well as methane emissions from plants and animals, all lead to a buildup of atmospheric greenhouse gases. Take asparagus: in a report prepared for the city of Seattle, Daniel J. Morgan of the University of Washington and his co-workers found that growing just half a pound of the vegetable in Peru emits greenhouse gases equivalent to 1.2 ounces of CO₂—as a result of applying insecticide and fertilizer, pumping water and running heavy, gas-guzzling farm equipment. To refrigerate and transport the vegetable to an American dinner table generates another two ounces of CO₂-equivalent greenhouse gases, for a total CO₂ equivalent of 3.2 ounces.

But that is nothing compared to beef. In 1999 Susan Subak, an ecological economist then at the University of East Anglia in England, found that, depending on the production method, cows emit between 2.5 and 4.7 ounces of methane for each pound of beef they produce. Because methane has roughly 23 times the global-warming potential of CO₂, those emissions are the equivalent of releasing between 3.6 and 6.8 pounds of CO₂ into the atmosphere for each pound of beef produced.

Raising animals also requires a large amount of feed per unit of body weight. In 2003 Lucas Reijnders of the University of Amsterdam and



Sam Soret of Loma Linda University estimated that producing a pound of beef protein for the table requires more than 10 pounds of plant protein—with all the emissions of greenhouse gases that grain farming entails. Finally, farms for raising animals produce numerous wastes that give rise to greenhouse gases.

Taking such factors into account, Subak calculated that producing a pound of beef in a feedlot, or concentrated animal feeding operation (CAFO) system, generates the equivalent of 14.8 pounds of CO₂—pound for pound, more than 36 times the CO₂-equivalent greenhouse gas emitted by producing asparagus. Even other common meats cannot match the impact of beef; I estimate that producing a pound of pork generates the equivalent of 3.8 pounds of CO₂; a pound of chicken generates 1.1 pounds of CO₂-equivalent greenhouse gases. And the economically efficient CAFO system, though certainly not the cleanest production method in terms of CO₂-equivalent greenhouse emissions, is far better than most: the FAO data I noted earlier imply that the world average emissions from producing a pound of beef are several times the CAFO amount.

Solutions?

What can be done? Improving waste management and farming practices would certainly reduce the “carbon footprint” of beef production. Methane-capturing systems, for instance, can put cows’ waste to use in generating electricity. But those systems remain too costly to be commercially viable.

Individuals, too, can reduce the effects of food production on planetary climate. To some degree, after all, our diets are a choice. By choosing more wisely, we can make a difference. Eating locally produced food, for instance, can reduce the need for transport—though food inefficiently shipped in small batches on trucks from nearby farms can turn out to save surprisingly little in greenhouse emissions. And in the U.S. and the rest of the developed world, people could eat less meat, particularly beef.

The graphics on the following pages quantify the links between beef production and greenhouse gases in sobering detail. The take-home lesson is clear: we ought to give careful thought to diet and its consequences for the planet if we are serious about limiting the emissions of greenhouse gases.

[THE AUTHOR]



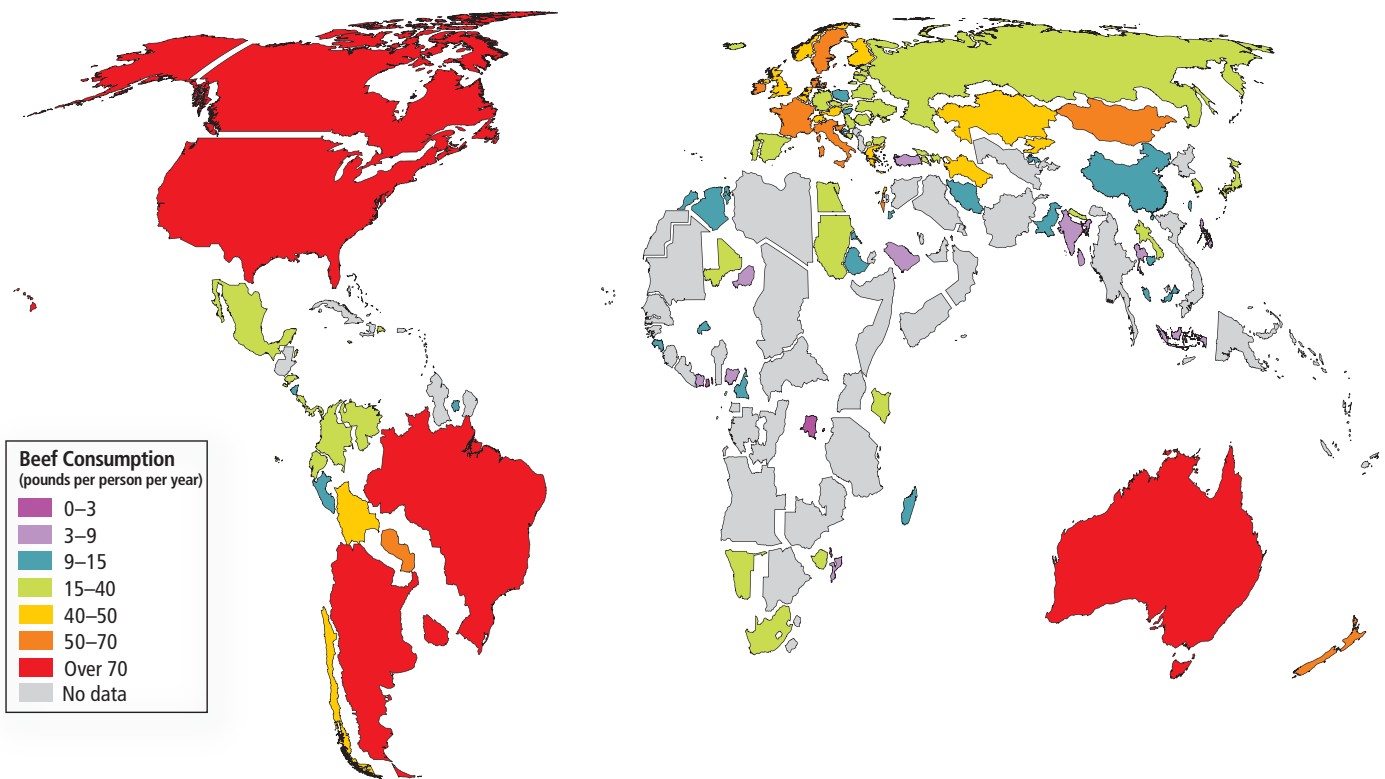
Nathan Fiala is a doctoral candidate in economics at the University of California, Irvine, focusing on the environmental impact of dietary habits. He also runs evaluations of development projects for the World Bank in Washington, D.C. In his spare time he enjoys independent movies and sailing. His study of the environmental impact of meat production on which this article is based was recently published in the journal *Ecological Economics*.

Burgers or Tofu?

Annual beef consumption per capita varies from 120 pounds in Argentina and 92 pounds in the U.S. to only a pound in the small eastern European country of Moldova; the average is about 22 pounds per person per year. The colors of the countries and the distortions of their usual shapes reflect the amount by which beef con-

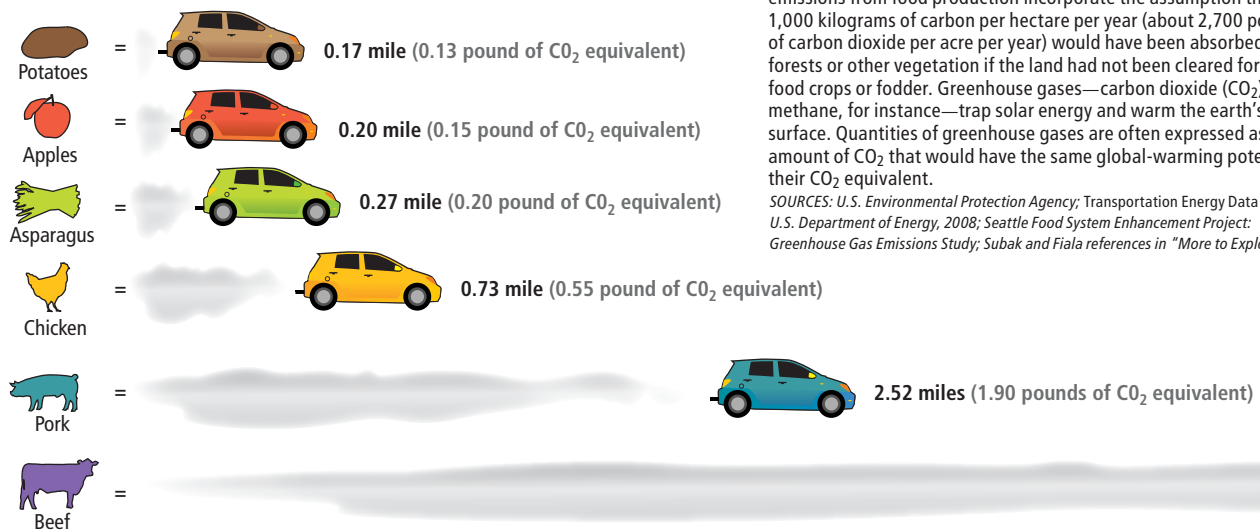
sumption per capita varies from the world average. World beef consumption per capita is growing, particularly in Asia, because of economic development: as people earn higher incomes, they purchase foods they find more desirable.

SOURCE: U.N. FAO, 2003



Eating and Driving: An Atmospheric Comparison

CO₂-equivalent emissions from producing half a pound of this food ... are the same as emissions from driving ...



The greenhouse gas emissions from producing various foods can be appreciated by comparing them with the emissions from a gasoline-powered passenger car that gets 27 miles per gallon. The estimated emissions from food production incorporate the assumption that 1,000 kilograms of carbon per hectare per year (about 2,700 pounds of carbon dioxide per acre per year) would have been absorbed by forests or other vegetation if the land had not been cleared for annual food crops or fodder. Greenhouse gases—carbon dioxide (CO₂) and methane, for instance—trap solar energy and warm the earth's surface. Quantities of greenhouse gases are often expressed as the amount of CO₂ that would have the same global-warming potential: their CO₂ equivalent.

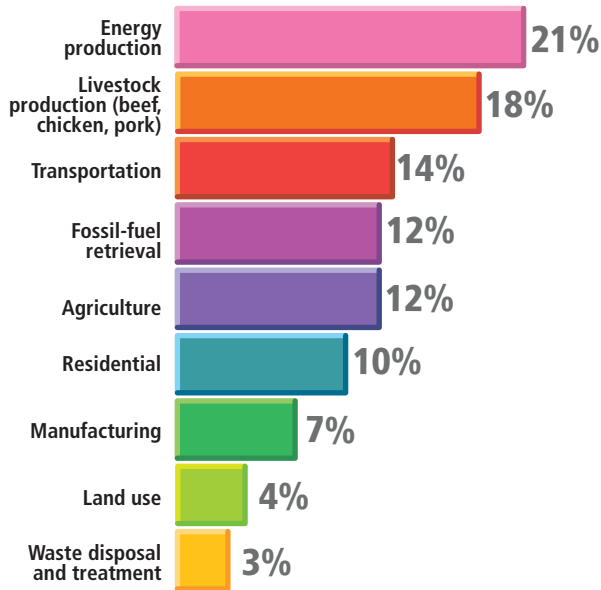
SOURCES: U.S. Environmental Protection Agency; Transportation Energy Data Book, U.S. Department of Energy, 2008; Seattle Food System Enhancement Project; Greenhouse Gas Emissions Study; Subak and Fiala references in "More to Explore"

MAPPING SPECIALISTS (map); LUCY READING-IKKANDA (illustrations)

The High (Greenhouse Gas) Cost of Meat

Worldwide meat production (beef, chicken and pork) emits more atmospheric greenhouse gases than do all forms of global transportation or industrial processes. On the basis of data from the United Nations Food and Agriculture Organization (FAO) and the Emission Database for Global Atmospheric Research, the author estimates that current levels of meat production add nearly 6.5 billion tons of CO₂-equivalent greenhouse gases every year to the atmosphere: some 18 percent of the worldwide annual production of 36 billion tons. Only energy production generates more greenhouse gases than does raising livestock for food.

SOURCE: U.N. FAO, 2006



Total is greater than 100% because of rounding



A Growing Appetite

	U.S. Beef Consumption (millions of tons)	CO ₂ -Equivalent Greenhouse Gases from U.S. Beef Production (millions of tons)	World Beef Consumption (millions of tons)	CO ₂ -Equivalent Greenhouse Gases from World Beef Production (millions of tons)
2009 (projected)	14	210	72	1,100
2020 (projected)	15	230	80	1,200
2030 (projected)	17	250	87	1,300
Cumulative (2009–2030)	340	5,000	1,800	26,000

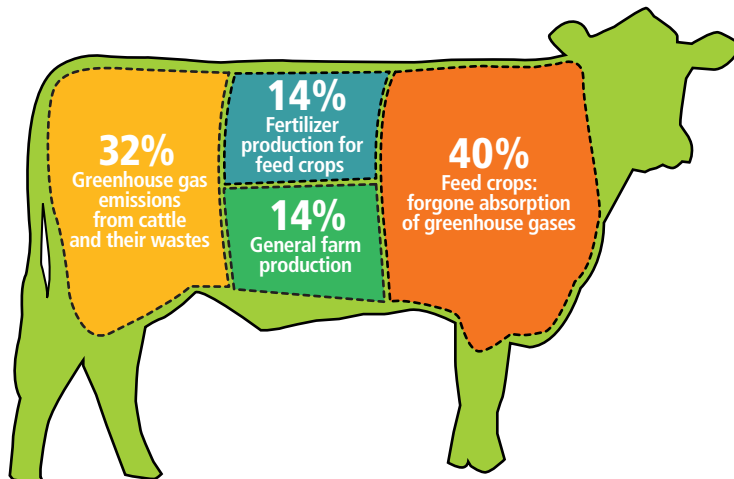
Quantities are rounded to two significant digits

World beef production is increasing at a rate of about 1 percent a year, in part because of population growth but also because of greater per capita demand in many countries. Economic analysis shows that if all beef were produced under the economically efficient feedlot, or CAFO (concentrated animal feeding operation), system—which generates fewer greenhouse emissions than many other common husbandry systems do—beef production by 2030 would still release 1.3 billion tons of CO₂-equivalent greenhouse gases. If current projections of beef consumption are correct, even under the feedlot production system the buildup of CO₂-equivalent greenhouse gases could amount to 26 billion tons in the next 21 years.

SOURCES: U.N. FAO; U.S. Census Bureau

Prime Cuts: How Beef Production Leads to Greenhouse Gases

The largest fraction of the greenhouse effect from beef production comes from the loss of CO₂-absorbing trees, grasses and other year-round plant cover on land where the feed crops are grown and harvested. Second most important is the methane given off by animal waste and by the animals themselves as they digest their food. This analysis of the U.S. feedlot beef production system was done by ecological economist Susan Subak, then at the University of East Anglia in England.



9.81 miles (7.40 pounds of CO₂ equivalent)



MORE TO EXPLORE

Global Environmental Costs of Beef Production. Susan Subak in *Ecological Economics*, Vol. 30, No. 1, pages 79–91; 1999.

Livestock's Long Shadow: Environmental Issues and Options. H. Steinfeld, P. Gerber, T. Wassenaar, V. Castel, M. Rosales and C. de Haan. United Nations Food and Agriculture Organization, 2006. Available at www.fao.org/docrep/010/a0701e/a0701e00.HTM

Meeting the Demand: An Estimation of Potential Future Greenhouse Gas Emissions from Meat Production. Nathan Fiala in *Ecological Economics*, Vol. 67, No. 3, pages 412–419; 2008.

Chaos and the Catch of the Day

There are fewer fish in the sea than ever. Complexity theory, argues mathematician George Sugihara, provides a counterintuitive way to revitalize the world's fisheries **BY PAUL RAEURN**

When George Sugihara reads about credit crises and federal bailouts, he is inclined to think about sardines—California sardines, to be precise.

A few decades after the Great Depression, the sardine fishery in California was suffering from a similarly devastating collapse. Fishers who had generally landed more than 500,000 tons of sardines annually during the 1930s caught fewer than 5,000 tons during the worst years of the 1950s and 1960s. Whereas a few Cassandras might have warned of trouble in each case, nobody could have predicted exactly when each collapse would come or how severe it would be.

The sardine collapse puzzled fisheries experts. Some blamed overfishing. Others suspected environmental swings—shifting wind patterns or cooling sea-surface temperatures. But nobody could prove either case. Eager to prevent another such collapse, California set up a monitoring system that has been collecting data on sardine larvae for the past 50 years.

Sugihara, a mathematician and theoretical ecologist at the Scripps Institution of Oceanography in La Jolla, Calif., analyzed that data and came to a surprising conclusion: both potential explanations of the sardine collapse were wrong.

His conclusion, in a study published in *Nature* in 2006,

was that the problem was the harvesting of too many big fish. Fishing boats were leaving behind a population of almost all juveniles. Sugihara showed that mathematically such populations are unstable. A slight nudge can create a boom—or a catastrophic collapse.

Imagine, Sugihara says, a 500-pound fish in an aquarium. Feed it more, and it gets fatter. Feed it less, and it gets thinner. The population (of one) is stable. But put 1,000

half-pound fish in that aquarium, and food shortages could result in the deaths of hundreds, because the small fry have less stored body fat—and therefore cannot ride out a short famine. Food abundance does not necessarily mean all the fish get bigger, either; it could encourage reproduction and a population boom—which might in turn overwhelm the food supply and lead to another bust. It is an unstable system. “That’s the reality of fisheries, of economies, of a

lot of natural systems,” Sugihara says. The recent history of the sardine fishery illustrates that instability: fishers along the West Coast from Canada to Mexico are now harvesting a million tons of sardines annually.

But this instability is not understood by people who run fisheries, Sugihara insists. By law they manage fisheries for “maximum yield.” The notion that such a maximum yield exists implies that fish grow at an equilibrium rate and that the harvest can be adjusted in accordance with that growth to keep yields stable. In contrast, Sugihara sees fisheries as a complex, chaotic system, akin to financial networks. They are so alike that the global financial giant Deutsche Bank lured Sugihara away from academia for a time; there, from 1996 to 2001, he successfully used the analytical techniques that he would later call on for his sardine work to make short-term predictions about market fluctuations.



GEORGE SUGIHARA

FOOD FOR THOUGHT: Using complexity theory, he has shown that standard fisheries practices produce unstable populations that can boom or bust even when food is abundant.

FISHY ADVICE: That fishers teach their children to throw the little ones back is exactly wrong when it comes to fisheries health, he says.

ON LIVING WITH CHAOS: “Most fisheries management is based on the idea that these systems are stable. Watches are like that. Transistors are like that. But ecosystems are not.”

Although both marine ecosystems and financial markets might look random, Sugihara explains, they are not. That means making short-term predictions is possible, as it is with the weather. The eminent ecologist Robert M. May of the University of Oxford calls that predictability “the flip side of chaos.” May oversaw Sugihara’s doctoral work at Princeton University when he was a visiting professor there and is now a frequent collaborator. “George was one of the first to see this as a recipe for making predictions,” he says.

Sugihara’s research comes at a time of enormous concern about the future of the world’s fisheries. Perhaps the most alarming report came in late 2006, when Boris Worm, a marine conservation ecologist at Dalhousie University in Nova Scotia, reported in *Science* that for 29 percent of currently fished species, the catch had dropped to less than 10 percent of the historical maximum. If the trends continue,

he reported, all fisheries around the globe will collapse by 2048.

Others think the future is not nearly so gloomy. “It’s very dependent on where you are,” comments Ray Hilborn, a professor of fisheries management at the University of Washington. The U.S., Canada and some other developed countries have cut fishing rates, and the future looks brighter, he says. But Asia and Africa lack effective fisheries management, and even European countries have failed to agree on solid management plans. Fisheries in those regions are in far greater peril, Hilborn states.

The practical implications of Sugihara’s work are clear. Current fishing regulations usually have minimum size limits to protect smaller fish. That, Sugihara maintains, is exactly wrong. “It’s not the young ones that should be thrown back but the larger, older fish that should be spared,” he explains. They stabilize the population and provide “more and better

quality offspring.” Laboratory experiments with captive fish back up Sugihara’s conclusions. For instance, David Conover of Stony Brook University found that harvesting larger Atlantic silversides from his tanks over five generations produced a population of smaller individuals.

Sugihara has also shown that populations of different fish species are linked. Most regulations consider each species—sardines, salmon or swordfish—in isolation. But fishing, he says, is like the stock market—the crash of one or two species, or a hedge fund or mortgage bank, can trigger a catastrophic collapse of the entire system.

Sugihara has also used his combined experience in ecology and finance to work on new kinds of fisheries management schemes. One is the notion of tradable “bycatch” credits. Bycatch refers to the turtles, sharks and other animals that fishing fleets do not seek but catch accidental-

EXERCISE IN EXACTLY 4 MINUTES PER DAY

How can this be true when experts insist that you need at least 30 minutes for cardio?

\$14,615



The so called “experts” have been wrong all these years. The length of a cardio workout depends on the level of oxygen consumption during the activity.

Walking consumes only 7 milliliters of oxygen per kg of bodyweight per minute (7 ml oxygen consumption / kg / min) and therefore must be done at least for 85 minutes to get an adequate cardio workout. Jogging produces 14 ml / kg / min and therefore only requires 35 minutes. Running at 6.5 miles per hour requires only 12 minutes and running at 4 minute mile speed requires only 3 minutes.

Our ROM machine is so designed that it will give anyone

from age 10 to 100 an excellent cardio workout in only 4 minutes. The next complaint from people is usually the \$14,615 price. It turns out that the ROM is the least expensive exercise by a wide margin.

Find out for yourself at www.PleaseDoTheMath.com. We have sold over 4800 ROM machines factory direct from our 82,000 sq. ft. factory in North Hollywood, CA.

Over 90% of our machines go to private homes, the rest to commercial uses such as gyms, doctor offices, military installations, and offices for the employees to use. Comes fully assembled, ready to use.

RENT A ROM FOR 30 DAYS. RENTAL APPLIES TO PURCHASE.

Order a **FREE DVD** from www.FastWorkout.com or call (818) 504-6450

PROUDLY MADE IN THE U.S.A. SINCE 1990

ROMFAB, 8137 Lankershim Blvd., North Hollywood, CA 91605 • Fax: (818) 301-0319 • Email: sales@FastWorkout.com

The typical ROM purchaser goes through several stages:

1. Total disbelief that the ROM can do all this in only 4 minutes.
2. Rhetorical (and sometimes hostile) questioning and ridicule.
3. Reading the ROM literature and reluctantly understanding it.
4. Taking a leap of faith and renting a ROM for 30 days.
5. Being highly impressed by the results and purchasing a ROM.
6. Becoming a ROM enthusiast and trying to persuade friends.
7. Being ignored and ridiculed by the friends who think you’ve lost your mind.
8. After a year of using the ROM your friends admiring your good shape.
9. You telling them (again) that you only exercise those 4 minutes per day.
10. Those friends reluctantly renting the ROM for a 30 day trial. Then the above cycle repeats from point 5 on down.

The more we tell people about the ROM the less they believe it.

From 4 minutes on the ROM you get the same results as from 20 to 45 minutes of aerobic exercise (jogging, running, etc.) for cardio and respiratory benefits, plus 45 minutes weight training for muscle tone and strength, plus 20 minutes stretching exercise for limberness and flexibility.

ly. In the tradable bycatch credits plan, fishing boats could be allocated a certain number of credits. As they used those credits, they would need to stop fishing or buy more credits on the open market. As the bycatch increased, the number of outstanding credits would fall, and their price would increase. Fishing boats would thus have financial incentive to minimize their bycatch—because by doing so, they could keep fishing longer.

Sugihara's work on fisheries has not met with universal acceptance. Roger Hewitt, assistant director of the National Oceanic and Atmospheric Administration's Southwest Fisheries Science Center in La Jolla, remarks that Sugihara's work is "a bit disconcerting" to people in fisheries management. "In fisheries, the classical approach is to model populations based on first principles. We know how fast [individual fish] are growing, how fast they are reproducing, how old they are when they mature,



NET OUTPUT: Sardine fishing sees good times and bad. Throwing the big ones back may help tame this unstable, complex system.

how many babies they have," Hewitt explains. "George's approach is an entirely different one. He looks at past behavior to see if he can predict future behavior." In a crude sense, Sugihara does not need to know about growth rates, reproduction or mortality.

Barry Gold, leader of the marine conservation effort at the Gordon and Betty Moore Foundation in San Francisco, de-

scribes Sugihara's analytical tools as "important for understanding how we manage fisheries." But he thinks that Sugihara's analysis needs a real-world test: "Until it's in the field and we see how the fishing industry responds to it, we won't know how it's going to work."

Partly in response to Gold and others about the lack of convincing field tests, Sugihara is now negotiating with fishing industry groups to try to put his work into practice. "Once you've stopped imagining that the world is a watch, that it's extremely predictable, you can make relatively good short-term forecasts," he states. "I have a lot of faith in human ingenuity. But the first step is acknowledging the reality."

Part of a 12-step program for fisheries ecologists? "Yes," Sugihara says. "It feels a little bit like that."

Paul Raeburn is a freelance science writer based in New York City.

JONATHAN S. BLAIR/National Geographic/Getty Images

Need Auto & Body Parts at Rockin' Prices?

Price Comparison

1994 TOYOTA PICKUP SR5 Timing Belt Component Kit

Parts Store	Part Brand	Price
RockAuto	Gates	\$127.79
O'Reilly	Gates	\$142.99
Autozone	CRP	\$179.99
Advance	Gates	\$197.99
Checker, Schucks, & Kragen	Gates	\$197.99
NAPA	NAPA	\$204.00

(as of 12/11/2008) RockAuto price is regular price, NOT a special sale price created for this comparison.

✓ Huge Selection
✓ Everyday Low Prices

✓ Fast Shipping
✓ Easy to use Website

ALL THE PARTS YOUR CAR WILL EVER NEED

GO TO WWW.ROCKAUTO.COM 1-866-ROCKAUTO (762-5288) ROCKAUTO, LLC - MADISON, WI USA

SCIENTIFIC AMERICAN

Subscriber alert!

Scientific American has been made aware that some subscribers have received notifications/subscription offers from companies such as United Publishers Network, Global Publication Services, Publishers Access Service, Lake Shore Publishers Service, Publishers Consolidated, Platinum Publishing Service, Publishers Service Center and American Consumer Publishing Association. These are not authorized representatives/agents of Scientific American. Please forward any correspondence you may receive from these companies to:

Christian Dorbandt
Scientific American
415 Madison Ave.
New York, NY 10017



Ski Butternut

**You could
drive further
and pay more...
but why?**



- * 22 trails
- * 12 lifts (3 quads)
- * 2 terrain parks
- * 1000' vertical
- * 100% snowmaking
- * 5 lane tubing center
- * 2 1/2 hrs from NYC or Boston

www.SkiButternut.com → Steals & Deals

\$20 **
Lift Tickets
* **MON-THURS**

Excludes holiday periods.

**midweek
ski & stay**

Packages from **\$45/day**
include lift ticket & lodging

* **Ski March
Free!**

When you buy your 09/10 Season Pass early:
\$229 adults **\$179** jrs **\$79** kids

The Berkshires' Affordable Family Mountain! **GREAT BARRINGTON, MA**

Visit our website for all our **GREAT DEALS** – Plus **SPECIALS** for college students, police, fire, EMS & more

Touch Screens Redefine the Market

By Mark Fischetti

In 2007, when Apple released the iPhone, its big touch screen made it an instant hit. The phone operated exclusively on AT&T's wireless network in the U.S., and other network providers implored their phone makers to quickly devise competitors. The scramble was on, and the touch-screen alternatives blossomed during the 2008 holiday season. Suddenly available were Research in Motion's BlackBerry Storm, which operates over Verizon's network, HTC's G1 (T-Mobile), the Samsung Instinct (Sprint), and others—most retailing for about \$200.

Each of these offerings can be called a smart phone, which generally means the technology is robust enough to provide a range of services beyond cell phone calls and text messaging and often means the operating system is open to third-party software developers seeking to create more novel features. The smart phones increasingly communicate over so-called 3G cellular networks that allow faster Web browsing and sending and receiving of e-mail. But the touch screens are the primary consumer draw. "Every provider now has a showcase phone that it is promoting heavily, to try to compete with the iPhone," says Ross Rubin, director of industry analysis at NPD Group, a market research firm in Port Washington, N.Y.

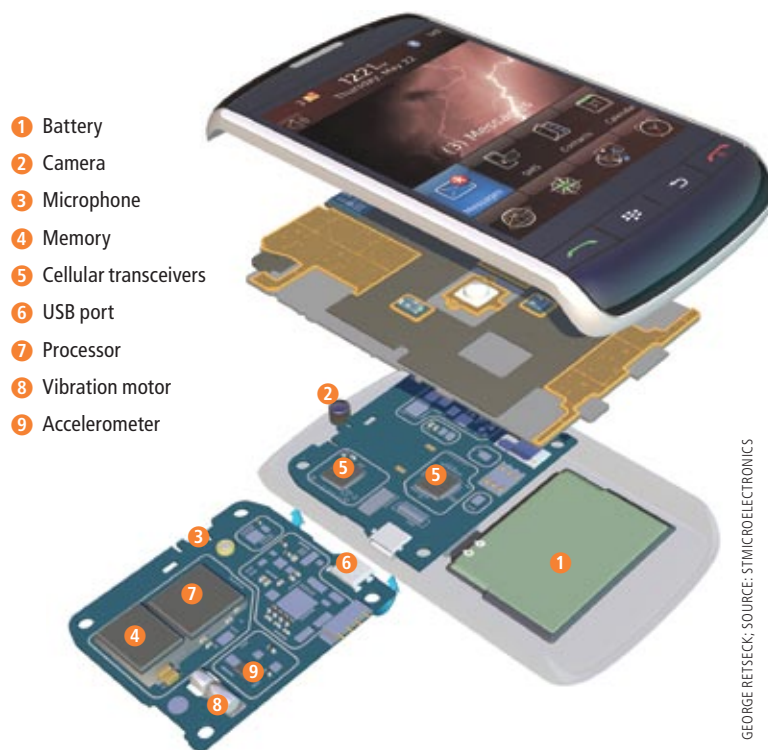
The handsets are crammed with hardware such as digital photograph and video cameras, music players and the dandy screens [see illustrations]. The devices may soon evolve into small computers about the size of a clutch purse. Hewlett-Packard and others have begun offering such "net-books" with 3G Internet capabilities; cell phone service is expected soon.

Smart phones pack an incredible array of telecommunications capabilities, including e-mail messengers, Web browsers, GPS navigators and, oh, yes, the actual cell phone. And more is to come. "There's plenty of gas left in the 3G network," Rubin points out. Eventually, though, the promise of delivering mobile broadband—equivalent to the DSL or cable service most people now enjoy at home—will require an evolution to 4G, which is already being planned under the monikers of LTE (AT&T and Verizon) and WiMAX (Sprint).

When 4G rolls out, carriers might finally open their wireless networks, so consumers could buy phones from different manufacturers that operate on various networks, whether AT&T's or Verizon's. The phones would probably be more expensive, because the carriers would not be subsidizing them to lock consumers into a two-year service contract. "You would just pay for monthly or even daily service," Rubin predicts, "with no penalty for switching."



➔ **SMART PHONES** such as the Apple iPhone (*above*) and BlackBerry Storm (*below*) are jammed with stand-alone components and telecommunications hardware.

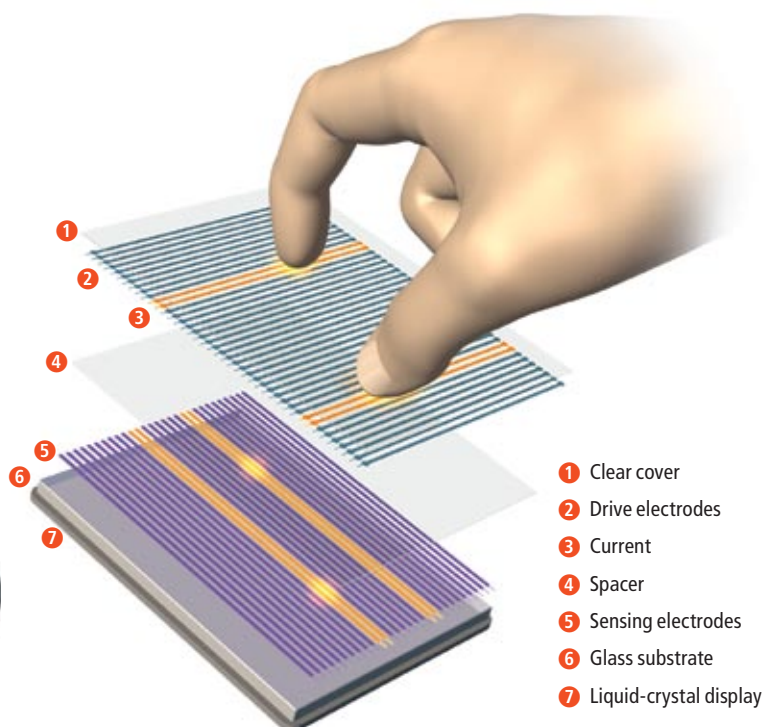


GEORGE RETSECK; SOURCE: STMICROELECTRONICS

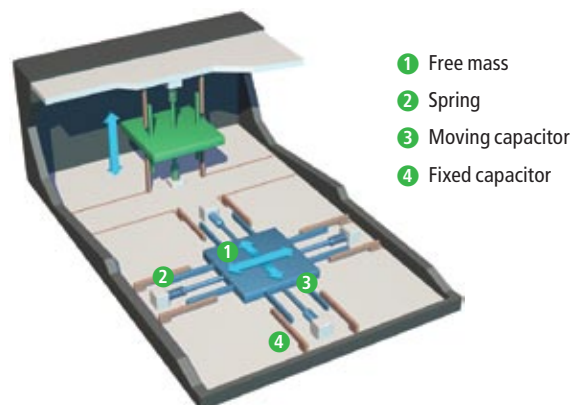
DID YOU KNOW ...

SHAKE IT: Accelerometers have been embedded in touch-screen phones to track when users turn the screens from "portrait" to "landscape" orientation. But their inclusion is allowing new applications, too. When the iPhone displays a list of nearby restaurants, shaking the phone will reorder the entries; the accelerometer can also instruct the camera to take a photograph at night in low light only if the phone is being held steady.

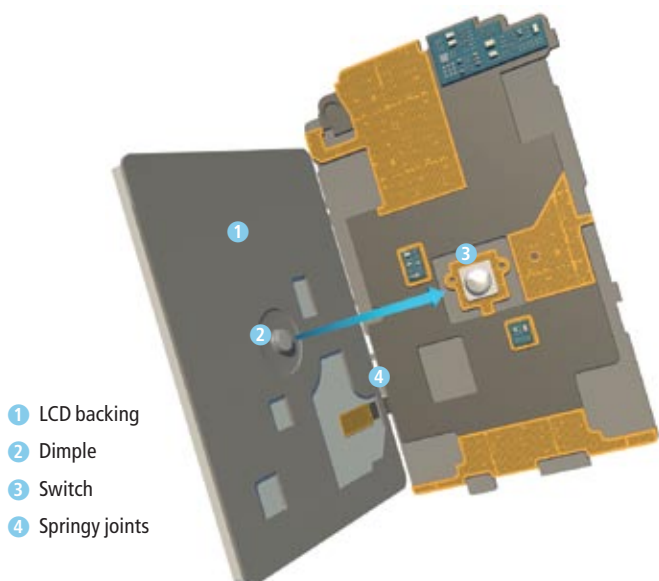
MIDS: Mobile Internet devices, or MIDs, are garnering more interest, as touch screens make cell phones bigger instead of smaller. MIDs are perhaps twice the size of touch-screen phones and are optimized for one function, such as a video camera that can wirelessly upload to the Web, or a mobile video-game console that allows people roaming around the world to play against one another. Intel Corporation is pushing the MID concept and name in part because the company makes a processor called Atom that can drive such devices and is already in very small "netbook" computers optimized for Web browsing.



➔ **SCREEN** in the iPhone and Storm is a "projected, mutual-capacitance" touch screen. Drive electrodes carry battery current and cross over sensing electrodes. When a conductive object such as a fingertip touches the screen, it alters the mutual capacitance across neighboring sensing electrodes, defining the point of contact. Most screens handle one touch at a time, but the iPhone multi-touch screen can respond to two fingers simultaneously.



➔ **ACCELEROMETER** senses when someone turns a phone from a vertical to horizontal orientation, so software can reshape imagery to fill the screen. In STMicroelectronics's three-axis, microelectromechanical design, when a free mass (*blue*) held by springs moves, attached capacitor plates pass by fixed plates, divulging the direction and extent of motion in the x and y planes. A second sensor (*green*) on the same chip tracks movement in the z direction.



➔ **CLICK SENSATION** in the Storm gives users tactile feedback when they press a virtual key on the screen. The entire suspended screen depresses slightly, and a dimple on the back side impinges on a microswitch. The switch pushes back up in response.

SEND TOPIC IDEAS to workingknowledge@SciAm.com

The Fourth Dimension ■ Extraterrestrial Life ■ Happy Birthday, Charles Darwin

BY MICHELLE PRESS

➔ IN SEARCH OF TIME: THE SCIENCE OF A CURIOUS DIMENSION

by Dan Falk. St. Martin's Press, 2008 (\$25.95)



First it was tools, then language. Now it is time that sets humans apart in the animal kingdom—the capacity to comprehend time, that is.

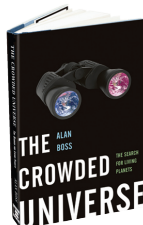
Science journalist Dan Falk

tackles the subject from several directions, starting with time's natural cycles (the heavens, the seasons) and moving on to the calendar; hours, minutes and seconds; memory; spacetime; time travel; illusion and reality. He hooks us from the outset, as he explores the prehistoric "passage tomb" at Newgrange in Ireland, more ancient than Stonehenge and much less known. There "a sliver of sunlight" briefly creeps into the dark chamber on the

morning of the winter solstice, allowing us "to glimpse, however dimly, into the minds of those who first considered the matter of time." From that point forward, Falk selects, organizes and interprets a mass of lore for our enlightenment and pleasure. We owe him.

➔ THE CROWDED UNIVERSE: THE SEARCH FOR LIVING PLANETS

by Alan Boss. Basic Books, 2009 (\$26)



In early March, NASA is scheduled to launch the Kepler Mission, the first space telescope specifically designed to detect habitable worlds orbiting stars similar to our sun. In point of fact, the Europeans were there first; in late 2006 they sent into orbit CoRoT, a space telescope designed to study stars but capable of detecting Earth-like

planets. Both missions have work ahead. The frequency of "habitable worlds," a rocky planet with liquid water near the surface where organisms can originate and evolve, is the most important factor in any estimate of the extent to which life has proliferated in the universe. Astronomer Alan Boss of the Carnegie Institution of Washington predicts that CoRoT and Kepler will discover abundant Earths. These telescopes are poised to prove him right or wrong, and his book provides essential and fascinating background as the drama unfolds.

NOTABLE BOOKS: DARWINIANA

Darwin's 200th birthday on February 12, 2009, together with the 150th anniversary of the publication of *On the Origin of Species*, has inspired a roster of books:

1 The Young Charles Darwin

by Keith Thomson. Yale University Press, 2009 (\$28)

Thomson, who manages to be stylish, scholarly and entertaining all at the same time, investigates Darwin's early years and how he arrived at his revolutionary ideas.

2 Darwin's Universe: Evolution from A to Z

by Richard Milner. University of California Press, 2009 (\$39.95)

All things Darwin, authoritative, amusing, abundantly illustrated, including some rare finds.

3 Remarkable Creatures: Epic Adventures in the Search for the Origin of Species

by Sean B. Carroll. Houghton Mifflin Harcourt, 2009 (\$26)

These "remarkable creatures" are in fact the pioneering naturalists whose dramatic discoveries fleshed out Darwin's theory. Carroll gives us voyages, expeditions, obstacles and what drives the passions of scientists.

4 On the Origin of Species: The Illustrated Edition

by Charles Darwin. Edited by David Quammen. Sterling, 2008 (\$35)

The most beautiful in the slate of Darwin books: over 300 paintings, photos, engravings and facsimiles.

5 Evolution: The First Four Billion Years

edited by Michael Ruse and Joseph Travis. Foreword by E. O. Wilson. Harvard University Press, 2009 (\$39.95)

The heaviest of the lot, at 1,008 pages, divided into a series of essays and an encyclopedic section that covers the spectrum of topics in evolution.



EXCERPT.....

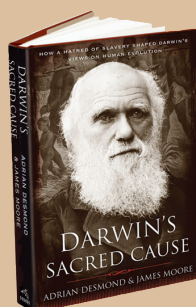
➔ DARWIN'S SACRED CAUSE: HOW A HATRED OF SLAVERY SHAPED DARWIN'S VIEWS OF HUMAN EVOLUTION

by Adrian Desmond and James Moore. Houghton Mifflin Harcourt, 2009 (\$30)

In this controversial reinterpretation of Charles Darwin's life and work, the authors of a highly regarded 1991 biography argue that the driving force behind Darwin's theory of evolution was his fierce abolitionism, which had deep family roots and was reinforced by his voyage on the Beagle and by events in America:

"... the barbarity of slavery brought his barely visible abolitionism into sudden sharp relief. It demanded a new commitment, the sort he had been unable to give till now.

"His encounters with Fuegians in cravats, and 'Hottentots' in white gloves, had proved human cultural adaptability as no anti-slavery tract could. Nothing could have better pointed up the pliancy of race. It justified the abolitionist faith in blacks being able to pull themselves up by their bootstraps. The impeccably mannered 'Hottentot' was living proof of the evil of considering 'wild' humans grovelling beasts. These racial encounters sharpened Darwin's sense of injustice—the tortured slaves, genocide of the Pampas Indians, the dragnet round-ups of the Tasmanians, the stuffing of 'Hottentot' skins. The result was an astonishing outpouring early in his notebooks which reveals the anguish behind his evolutionary venture."



Reprinted by permission of Houghton Mifflin Harcourt Publishing Co. All rights reserved.

SciAm Marketplace

SciAm.com

Complement
your print presence
with an
online ad on
SciAm.com



Contact

Beth Boyle at

914.461.3269

or beth@specialaditions.com

for more information

Portable FUN In the SUN!



Our Great NEW Kayak Goes Anywhere!

Our 12' 4" Sea Eagle 370 inflatable kayak holds 2 adults, weighs just 34 lbs. & packs in small nylon bag, yet it will carry up to 650 lbs. of people & gear. Paddle wild rivers, remote ponds, secluded bays, scenic lakes... even ocean surf!

SE-370 Pro Package includes 2 7'10" aluminum 4-part backpack paddles, 2 deluxe kayak seats, foot pump, nylon carry bag, instructions and repair kit. All fit in the convenient carry bag.
NOW ONLY \$399. COMPLETE - Save \$184!

But that's not all, this winter you also save the \$39. shipping charge because **SHIPPING IS FREE!**
Also get 6 Month Money Back Trial Guarantee & 3 Year Warranty Against Manufacturing Defects!

SEA EAGLE.com Dept SC029B

Visit **SeaEagle.com** to order & for more info
or Call for Catalog **1-800-748-8066**
19 N. Columbia St., Port Jefferson, NY 11777

Nova Scotia

\$995 8 Day Tour
Call for Choice Dates

Choose Your Tour—Only \$995

- 9 days Canadian Rockies
- 8 days Nova Scotia, Prince Edward Isl.
- 8 days Grand Canyon, Zion, Bryce
- 8 days Mount Rushmore, Yellowstone
- 8 days Mexico's Copper Canyon Train
- 11 days Guatemala
- 10 days Costa Rica

FREE
Info
Guide

Fully Escorted Tours Since 1952
Caravan
1-800-Caravan Caravan.com

Most Caravan tours are \$995 pp, dbl. occ., plus taxes and fees



Save thousands of dollars!
Do more with geographic data.

Manifold® delivers the world's most powerful, most complete, easiest to use GIS at the world's lowest price. From \$245. Enterprise x64 for \$445. Web apps too! Supercomputer speed and total reliability.
www.manifold.net

Whole House Fan
COOL YOUR HOME
3 Steps To A Cooler Home:

1. Cool fresh evening air is drawn into the home via open windows, screen doors, etc.
2. Warm stale indoor air is vented into the attic space with a WH Whole House Fan.
3. Superheated attic air is exhausted out of the home - WHO GET COOL - Good Night!

WholeHouseFan.com
Stay Cool & Save \$\$\$
Find Out How...
WholeHouseFan.com
Toll Free: 888-845-6597

Makes Glass INVISIBLE

See better than you've ever imagined. Improve your driving visibility. Invisible Glass contains a powerful, non-abrasive formula with NO streaky soaps or foams, for optimum clarity. Makes windshields, windows and mirrors disappear. Trial offers, more information at www.invisibleglass.com or 1-888-4STONER, M-F, 8-5pm. Ad Code VSA9A



EXERCISE IN EXACTLY 4 MINUTES PER DAY

\$14,615



Most people think that the 4 minutes and the \$14,615 price must be mistakes. They are not. Since 1990 we have shipped over 4800 ROM machines from our 82,000 sq. ft. factory in North Hollywood, CA. People first rent it for 30 days at \$2500 to try it. Rental applies to purchase. Then 97% buy it.

Find out why this is the least expensive exercise
www.PleaseDoTheMath.com

Order a **FREE DVD**
from **www.FastWorkout.com**
or call (818) 504-6450



Why do wind turbines have three narrow blades, whereas my fan at home has five wide blades?

—J. Lester, Stroudsburg, Pa.

Dale E. Berg, a member of the technical staff in the wind energy technology department at Sandia National Laboratories, replies:

The differences between wind turbine and ceiling fan blades arise from the contrasting design criteria: the wind turbine is intended to capture high-velocity wind to generate electricity efficiently; the ceiling fan needs to move air at low velocity with inexpensive components.

To keep drivetrain costs low, a wind turbine must capture the energy in fast-moving air and rotate at relatively high speed—within limits, so as to avoid excessive noise generation. (Slow rotation would increase the torque and require heavier and more expensive drivetrain components.) Such high-efficiency energy conversion dictates the use of lift-type turbine blades, similar to airplane wings, of twisted and tapered airfoil shapes. The blade design creates a pressure difference in wind—high pressure on one side and low pressure on the other—that causes the blades to turn. A combination of structural and economic considerations drives the use of three slender blades on most wind turbines—using one or two blades means more complex structural dynamics, and more blades means greater expense for the blades and the blade attachments to the turbine.

The ceiling fan, on the other hand, is built to keep the occupants of a room comfortable by moving air gently. Its engineers work to minimize noise while the fan rotates at low speed (for safety reasons) and to keep the construction costs, and therefore the purchase price, low. Energy efficiency is not a primary concern, because operation is inexpensive—a typical ceiling fan running 24 hours a day consumes about 60 kilowatt-hours a month, for an average electricity cost of about six dollars. For this reason, most ceiling fans incorporate blades that are comparatively inefficient drag devices; rotating the pitched blades pushes air vertically out of the way. Wide, flat blades are inexpensive to build and work well as drag devices. More blades are better, up to a point, and the usual layout of four or five blades is the result of balancing trade-offs between efficiency and expense.

A 2001 article in *Mechanical Engineering* chronicled the quest of a man named Danny Parker to create a more ef-

ficient ceiling fan. Parker's initial blade prototype looked a lot like a wind turbine blade, but the end result (because of manufacturing, safety and operating concerns) was a hybrid between a standard ceiling fan blade and a wind turbine blade.

What happens to the donor's DNA in a blood transfusion?

—W. McFarland, Winter Springs, Fla.

Michelle N. Gong, an assistant professor at the Mount Sinai School of Medicine, explains:

Studies have shown that donor DNA in blood transfusion recipients persists for a number of days, sometimes longer, but its presence is unlikely to alter genetic tests significantly. Red blood cells, the primary component in transfusions, have no nucleus and no DNA. Transfused blood does, however, host a significant amount of DNA-containing white blood cells, or leukocytes—around a billion cells per unit (roughly one pint) of blood. Even blood components that have been filtered to remove donor white cells can have millions of leukocytes per unit.

Investigators have detected donor DNA after transfusion with a process called polymerase chain reaction (PCR) that amplifies minuscule amounts of genetic material for detection and identification of specific genes.

Studies using PCR to amplify male genes in female recipients of transfusions from male donors have demonstrated that donor DNA endures in recipients for up to seven days. And a study of female trauma patients receiving large transfusions showed the presence of donor leukocytes for up to a year and a half.

All these results, however, were found using very sensitive techniques whereby donor DNA was selectively amplified over the more plentiful recipient DNA. In studies where genes common to both donors and recipients were amplified, the results reflected the dominance of the transfusion recipient's own DNA, showing the donor's DNA to be a relatively inconsequential interloper.

HAVE A QUESTION?... Send it to
experts@SciAm.com or go to
www.SciAm.com/asktheexperts



Panama Canal / Southern Caribbean

February 27 – March 9, 2009

Evolution Emanation™

www.InSightCruises.com/SciAm3

Travel the Scientific American Way!
Bright Horizons™

Are you curious about the evolution of "Evolution"?

Get the new angles on this important scientific discipline and topic of debate and policy. Journey with Scientific American into the Panama Canal whose engineering is the product of an historic struggle against Nature and skepticism. Sail with kindred spirits on SciAm's Evolution Emanation cruise-conference on Holland America's Zuiderdam round-trip Fort Lauderdale.

Find out the relationships between key pieces of logic, methods, and facts documenting evolution. See beyond the obvious into studies of the genetic code, the nuances of speciation and natural selection, and a meaning of immortality. Let yourself be absorbed in discussion of the impacts of cognitive psychology on philosophy, and of the emerging mathematics of the mind.

Cruise prices vary from \$1,529 for a Better Inside to \$5,199 for a Full Suite, per person. (Cruise pricing is subject to change. InSight Cruises will generally match the cruise pricing offered at the Holland America website at the time of booking.) For those attending the conference, there is a \$1,275 fee. Taxes, gratuities, and a fuel surcharge are \$405.

While you're in balmy tropical climes, why not kayak with a companion through crystal blue waters, stride through Curaçao's unique nature sanctuaries, and venture into Costa Rica's rainforests, surrounded by the sounds of wildlife and colors of flora and fauna. Prioritize your explorations, or just let them flow . . .

Sail with Scientific American and relax and refresh while gearing up to a refined understanding of the science of evolution, the challenges it faces, and the benefits it brings. Visit www.InSightCruises.com or call 650-787-5665 to get all the details, and then enrich mind and body with an intellectual adventure with the Scientific American community.



Listed is a sampling of the
20+ sessions you can participate in
while we're at sea.
For a full listing visit
www.InSightCruises.com/SciAm3-talks

The Evolution of the Genetic Code

Speaker: Stephen J. Freeland, Ph.D.

Why should all life use two utterly different chemical languages with which to construct itself? Why did it arrive at a single set of encoding rules for translating between them? How can the precise choice of coding rules influence life's struggle to stay one step ahead of extinction? You'll get the bottom line (as it stands now) on the origin and subsequent evolution of the genetic code from Dr. Freeland, including:

- the central dogma that unifies life
- the non-random "design" of genetic code words
- emergence from an RNA world
- the great unknowns

1859: The Impact of a Dangerous Idea

Speaker: Jerry Coyne, Ph.D.

In this session we'll trace the origin of Darwin's "dangerous idea" (actually several ideas) beginning with his famous voyage on the HMS Beagle. We will learn what Darwin really proposed, what impact the ideas of evolution and natural selection had on the Victorian world, and why Darwinism was — and still is — considered a dangerous idea.

Unconscious Design: Natural Selection

Speaker: Jerry Coyne, Ph.D.

While the idea of evolution was immediately accepted by 19th-century biologists, the concept of natural selection — the purposeless driving force of evolution and adaptation — has been much more controversial. This talk will describe what natural selection really is and see examples of how it works in nature. We will also examine the complementary theory of sexual selection, which explains the remarkable difference in appearance and behaviour between males and females in many species.

For details contact:

Neil Bauman • 650-787-5665
or neil@InSightCruises.com

On the Origin of Species, Really

Speaker: Mohamed Noor, Ph.D.

Although Darwin's book title suggested that he provided us with insights on the origin of species, in fact, he only focused on the process of divergence within species and assumed the same processes "eventually" led to something that could be called a new species. In this session, we'll talk about how species are identified (in practice and in principle), and then how modern evolutionary biologists use this type of information to get a handle on how species are formed.

From Magic to Muons: Why People Believe in Strange Things

Speaker: Tania Lombrozo, Ph.D.

Much of our knowledge is about things that we cannot see or touch. By studying human reasoning we can begin to understand both how people make scientific discoveries and how these processes can lead to some surprising errors in understanding our world. We'll consider the debate over evolution and intelligent design as a case study in people's understanding of and preference for different kinds of explanations for the world around us.

The Mathematics of Mind: Exploring the Formal Foundations of Human Thought

Speaker: Thomas Griffiths, Ph.D.

Over the last two millennia, scientists and philosophers have used approaches such as logic, artificial neural networks, and probability theory to develop scientific and mathematical models of thought. Dr. Griffiths will talk about current status of work to understand the formal principles that underlie human thought and our ability to solve the computational problems we face in everyday life.

Evolution of Individuality and Complexity Through Cooperation and Conflict

Speaker: Richard Michod, Ph.D.

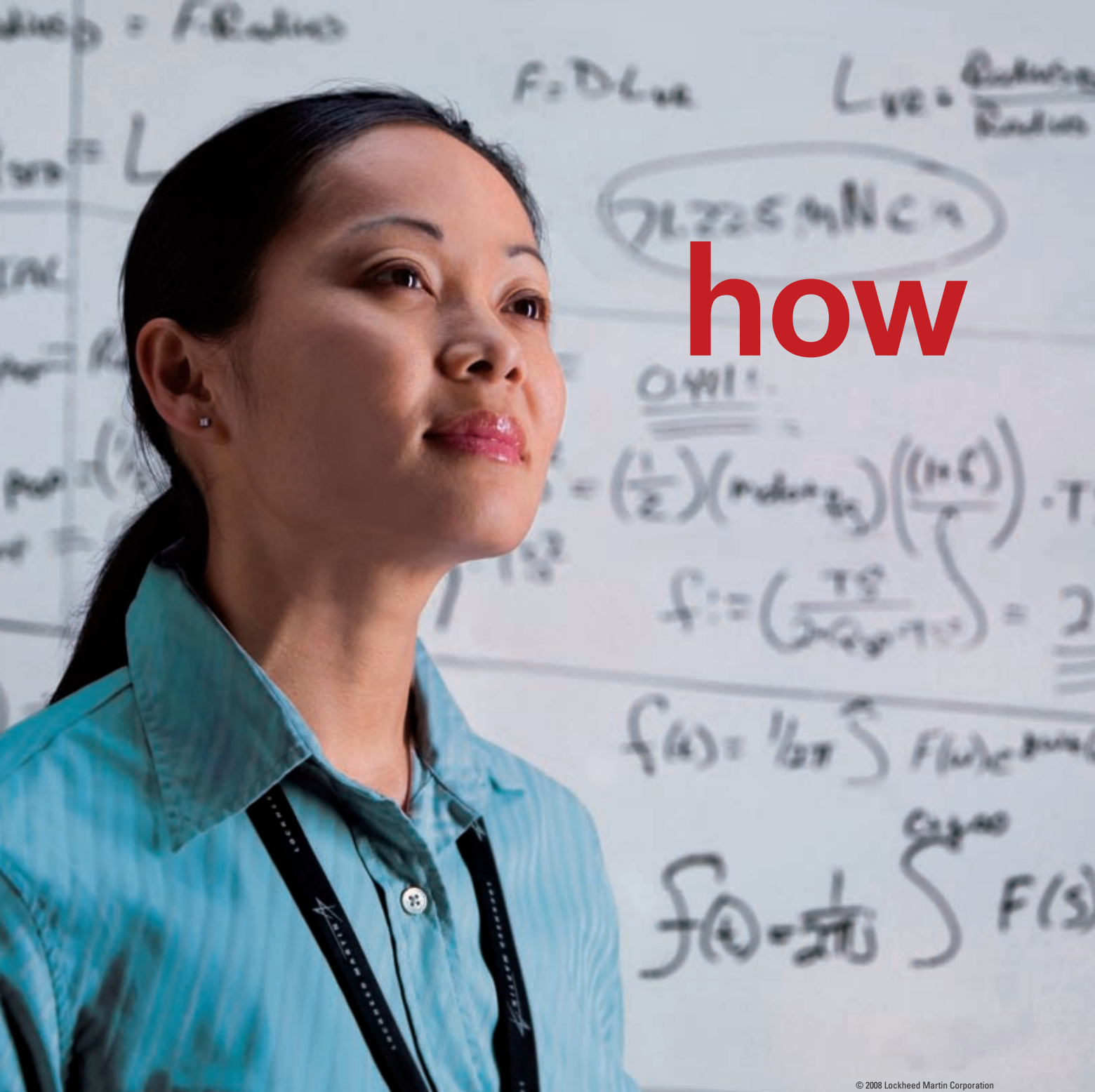
Our understanding of life is being transformed by the realization that evolution occurs not only among individuals within populations, but also through the integration of groups of cooperating individuals into new higher-level individuals — that is, through evolutionary transitions in individuality (ETIs). The major landmarks in the diversification of life and the hierarchical organization of the living world are consequences of a series of ETIs: from genes to gene networks to the first cell; from prokaryotic to eukaryotic cells; from cells to multicellular organisms; from asexually reproducing individuals to sexually reproducing pairs; and from solitary individuals to societies. How do groups become new individuals? Cooperation and conflict play a major role in these evolutionary transitions. Join Dr. Michod and come away with a new perspective on the process of evolution and what it means to be an individual.

BRIGHT HORIZONS CO-PRODUCED BY:

SCIENTIFIC
AMERICAN

InSight Cruises
EDUCATION THAT TAKES YOU PLACES

CSF# 2065380-40



how

BETWEEN THE IDEA AND THE ACHIEVEMENT,
THERE IS ONE IMPORTANT WORD: HOW.

AND IT IS THE HOW THAT MAKES ALL THE DIFFERENCE.

lockheedmartin.com/how

LOCKHEED MARTIN
We never forget who we're working for®